

Jean-Luc Montagnier

Construire son réseau d'entreprise



Au sommaire

- Choisir un accès Internet
- Construire un réseau local : Ethernet et IP
- Installer le câblage
- Dimensionner et configurer son réseau
- Interconnecter des sites distants
- Bâtir un réseau de transport : LS, Frame Relay ou ATM
- Voix et vidéo sur IP
- Administrer son réseau IP

VISIT...

LANZAROTE
Caliente.COM

Construire
ré^{son}seau
d'entreprise

CHEZ LE MÊME ÉDITEUR

Du même auteur

J.-L. MONTAGNIER. – **Pratique des réseaux d'entreprise.**
N°9031, 1998, 552 pages.

Ouvrages réseaux et télécoms

G. PUJOLLE. – **Initiation aux réseaux.**
N°9155, 2000, 448 pages.

G. PUJOLLE. – **Les réseaux.**
N°9119, 3^e édition, 2000, 950 pages.

J.-F. SUSBIELLE. – **L'Internet multimédia et temps réel.**
Réseaux haut débit – Terminaux fixes et mobiles – routage et QoS – voix et audio-vidéo sur IP.
N°9118, 2000, 750 pages.

J.-L. MÉLIN. – **Pratique des réseaux ATM.**
N°8970, 1997, 280 pages.

I. RUDENKO. – **Configuration IP des routeurs Cisco.**
N°9238, 2001, 386 pages.

C. LEWIS. – **Installer et configurer un routeur Cisco.**
N°9102, 1999, 450 pages.

Solutions Linux et Windows 2000

H. HOLZ, B. SCHMITT, A. TIKART. – **Internet et intranet sous Linux.**
N°9101, 1999, 474 pages + CD-Rom.

D.L. SHINDER, T. SHINDER – **TCP/IP sous Windows 2000.**
N°9219, 2001, 540 pages.

D.L. SHINDER, T. SHINDER – **Administrer les services réseau sous Windows 2000.**
N°9168, 2000, 600 pages.

M. CRAFT – **Active Directory pour Windows 2000 Server.**
N°9167, 2000, 360 pages.

Réseaux et services WAP

WAP FORUM – **Le guide officiel WAP1.2.**
N°9186, 2001, 1200 pages.

D. JAMOIS-DESAUTEL. – **Guide des services WAP.**
N°9257, 2001, 250 pages.

L. LETOURMY, T. PAPIERNIK, A. HELAÏLI, X. MARTZEL. – **Construire une application Wap.**
N°9174, 2000, 360 pages.

S. MANN. – **Initiation à WAPet WML.**
N°9179, 2000, 240 pages.

T. ZIEGLER. – **Précis WAPet WML.**
Les bases de la programmation. N°9249, 2000, 110 pages.

Jean-Luc Montagnier

Construire ré^{son}seau d'entreprise



ÉDITIONS EYROLLES
61, Bld Saint-Germain
75240 Paris Cedex 05
www.editions-eyrolles.com



Le code de la propriété intellectuelle du 1^{er} juillet 1992 interdit en effet expressément la photocopie à usage collectif sans autorisation des ayants droit. Or, cette pratique s'est généralisée notamment dans les établissements d'enseignement, provoquant une baisse brutale des achats de livres, au point que la possibilité même pour les auteurs de créer des œuvres nouvelles et de les faire éditer correctement est aujourd'hui menacée.

En application de la loi du 11 mars 1957, il est interdit de reproduire intégralement ou partiellement le présent ouvrage, sur quelque support que ce soit, sans autorisation de l'Éditeur ou du Centre Français d'Exploitation du Droit de Copie, 20, rue des Grands-Augustins, 75006 Paris.

© Éditions Eyrolles, 2001, ISBN 2-212-09262-8

ISBN édition Adobe eBook Reader : 2-212-28184-6, 2002

Distribution numérique par GiantChair, Inc.

Préface

Si le passage au troisième millénaire n'a heureusement pas été marqué par toutes les catastrophes que quelques prédicateurs affabulateurs en mal de reconnaissance avaient pourtant annoncées, il aura été à coup sûr le théâtre d'un profond bouleversement économique dont les historiens de demain sauront dire s'il aura constitué une révolution d'ampleur comparable à celle qu'a engendrée par exemple le développement de l'automobile, pour ne prendre qu'un exemple assez récent.

En effet, en à peine trente ans, une économie nouvelle aura vu le jour. Ses géniteurs, qui sont nés ou qui étaient adolescents pendant les années 1970-80, probablement marqués par la récession qui accompagna ces années, auront eu la géniale idée de fonder cette économie non pas sur l'exploitation d'une quelconque nouvelle richesse naturelle mais sur celle d'une matière première d'un tout nouveau genre, existant en quantité infinie, et de surcroît en perpétuelle création. Cette nouvelle richesse aux caractéristiques remarquables s'il en est, c'est tout simplement l'information. Qu'il s'agisse d'images, de voix ou de données, les métiers qu'elle a su développer visent à sa création, à sa préparation et avant tout à son partage.

Au cœur de cette nouvelle économie se trouvent bien sûr les réseaux de télécommunications sans lesquels rien n'aurait été possible. Qu'ils soient locaux, de campus, métropolitains, radio ou encore étendus, ils sont le complément nécessaire de l'information, son moyen de transport, et donc la condition indispensable de son partage. À tous les niveaux, ils consti-

tuent de véritables enjeux stratégiques. Dans l'entreprise, ils sont si présents qu'ils sont devenus indispensables à la réalisation de tous ses métiers ; dans la vie de tous les jours, ils occupent une place de plus en plus importante ; pour le monde industriel des constructeurs et opérateurs, ils sont l'objet de luttes féroces ; pour les États enfin, ils sont la source de revenus substantiels et sont de plus en plus partie prenante dans les missions de sécurité et de défense.

Mais si l'on peut être convaincu qu'*a posteriori*, on constatera que l'incroyable développement des technologies et que les investissements colossaux consentis auront permis la construction de réseaux de télécommunication toujours plus performants, plus fiables et plus accessibles aux usagers (condition *sine qua non* du succès de l'ensemble), force est de constater que le spectacle de leur réalisation dresse un tableau paradoxal, où s'entremêlent l'enthousiasme et l'euphorie des uns avec la désillusion et le désarroi des autres.

Prenons un exemple pour illustrer nos propos, celui des réseaux de communication sans fil. Ils répondent indéniablement à un véritable besoin de communication pour les usagers, et le succès des réseaux GSM ou DCS est incontestable. Certains projets ont pourtant d'ores et déjà connus des dénouements moins favorables. C'est le cas notamment des projets de réseaux de communication sans fil fondés sur l'utilisation d'un maillage de satellites comme infrastructure de raccordement et de transit.

Dans le même ordre d'idées, on peut dire que les espoirs placés dans l'émergence des réseaux à la norme UMTS peuvent être mesurés au travers des montants colossaux investis, d'une part, par les opérateurs de télécommunication pour l'attribution des licences d'exploitation et pour la construction des nouvelles infrastructures et, d'autre part, par les constructeurs pour le développement de cette nouvelle technologie. Mais à peine sommes-nous remis de nos émotions que le doute s'installe déjà, que certains élus des premières batailles revoient leur stratégie et se demandent si le succès sera aussi total que les premiers augures ne le laissaient croire.

Cet apparent paradoxe n'est pas une caractéristique singulière des réseaux de télécommunication. Une autre illustration nous en est donnée par les start-up nées avec la vague de l'Internet et qui ont largement contribué au bouleversement économique auquel nous nous intéressons. Hier encore, stars du Nasdaq attirant sans efforts les capitaux des investisseurs ainsi que les meilleurs ingénieurs avec le mirage des stocks options, aujourd'hui, et pour les meilleures uniquement, entreprises en phase de maturité à la recherche de la véritable création de valeur, faisant l'expérience souvent douloureuse du retour aux fondamentaux.

Ces quelques exemples, et il y en a bien d'autres, illustrent assez clairement que les évolutions des techniques et technologies ne se font jamais sans désillusion, remise en cause et même échec cinglant. Cela est en particulier vrai dans les domaines « à la mode » où les entreprises, fortes des moyens dont elles disposent de manière conjoncturelle, ne mettent pas en place les filtres amont qui permettraient de ne retenir que les avancées réellement viables. Il est donc du devoir de tout responsable télécommunication d'entreprise de s'assurer que les décisions qu'il prend en matière de choix de solution réseau sont bien le résultat d'une analyse approfondie où il aura évalué, entre autres, la maturité et la pérennité des solutions pro-

posées (et la crédibilité de ceux qui les proposent), leurs performances, leur fiabilité, leur coût de possession (et pas seulement leur coût d'acquisition) et la qualité du support proposé.

L'objet de cet ouvrage est de faire le point sur les techniques et technologies qui ont acquis un niveau de maturité nécessaire pour une mise en œuvre « dérisquée » ou qui sont en passe de l'acquérir. Il propose aux lecteurs à la fois les informations élémentaires nécessaires à leur compréhension et les clés de leur mise en œuvre. Son approche progressive des solutions et architectures réseau permettra aux lecteurs, en fonction de ses centres d'intérêts et de ses besoins, de trouver les réponses adaptées à ces interrogations.

Olivier Koczan

EADS Defence and Security Networks

Table des matières

CHAPITRE 1

Introduction	1
Le principe de l'apport minimal	2
Un exemple dans le domaine des réseaux	2
<i>A-t-on besoin d'une qualité de service ?</i>	3
<i>L'intégration voix/données/image</i>	3
<i>Les débits</i>	3
Le cas ATM	4
Pourquoi ATM pourrait s'imposer ?	4
Pourquoi ATM pourrait disparaître ?	4
ATM peut-il remplacer IP ?	5
Le syndrome de Vinci	5
Les technologies gagnantes	6
Les évolutions	7

PREMIÈRE PARTIE SE LANCER DANS LES RÉSEAUX

CHAPITRE 2

Premier contact avec les réseaux : l'Internet	11
Qu'est-ce que l'Internet ?	12
De quoi a-t-on besoin ?	12
D'un modem	12
D'un câble série	13
D'une prise téléphonique	14
D'un abonnement à l'Internet	14
De logiciels de communication	15
Comment faire fonctionner tout cela ?	17
Identifier le port série sur le PC et y connecter le modem	17
Installer et configurer le pilote du modem	18
Configurer l'accès réseau à distance	19
Tester votre connexion	22
Résoudre les problèmes de connexion	23

CHAPITRE 3

Les applications du réseau : l'exemple du web	25
Configurer et utiliser un navigateur	26
Qu'est-ce que le web ?	26
Le nommage des pages web	26
Naviguer	27
Améliorer les performances	29
Les cookies	31
Transférer des fichiers	32
Installer et configurer la messagerie	33
Les annuaires	38
Mais qu'est-ce que l'Internet ?	39

Quelques chiffres	41
Qui gère l'Internet ?	41
Quelques autres organismes d'intérêt général	44
Les anciens organismes de régulation	44
Où les contacter ?	45

CHAPITRE 4

Construire son premier réseau local	47
Le contexte	48
Les choix de base	49
Quel réseau ?	49
Quelle topologie ?	49
De quoi a-t-on besoin ?	51
De cartes réseau	51
De cordons de raccordement	53
D'un concentrateur	54
De logiciels de communications	57
Comment faire fonctionner tout cela ?	59
Installer les cartes réseau et les drivers	59
Configurer les adresses IP	61
Installer les concentrateurs et y raccorder les PC	63

CHAPITRE 5

Mettre en place un système de câblage	65
Quelle est la démarche à suivre ?	66
L'avant-projet	66
L'étude d'ingénierie	70
Quel type de câble ?	70
<i>Cuivre ou fibre optique ?</i>	71
<i>Coaxial ou paires torsadées ?</i>	72
Le choix de la paire torsadée en distribution	72
<i>Quelle impédance : 100, 120 ou 150 Ohms ?</i>	72

<i>Écranté ou non ? Blindé ou non ?</i>	72
<i>Catégories 5, 6 ou 7 ?</i>	74
Le choix de la fibre optique entre les locaux techniques	74
<i>Multimode ou monomode ?</i>	75
<i>62,5/125 ou 50/125 ?</i>	76
<i>Le câble contenant les fibres</i>	76
Le coaxial et la paire torsadée pour la vidéo	76
<i>Sur quels critères choisir le type de câble ?</i>	78
Quel type de prise ?	78
L'aménagement des locaux techniques	79
Les baies	79
Le cheminement des cordons de brassage	79
L'organisation du local	80
Le cahier des charges	80
Le suivi du chantier et la recette	83

CHAPITRE 6

Architecture des réseaux locaux	85
Les choix de base	86
Quel type de réseau choisir ?	86
Quel débit retenir ?	86
Quel format d'équipement ?	87
Concentrateur ou commutateur ?	89
L'architecture	92
Mise en place d'un réseau local d'étage	92
Extension du réseau d'étage	94
Conception d'un réseau d'immeuble	95
Mise en place d'un réseau fédérateur	96
Quel débit et quelle technologie ?	98
Suivre l'évolution des besoins	100
Assurer la continuité de service	102

DEUXIÈME PARTIE INTERCONNECTER SES RÉSEAUX

CHAPITRE 7

Démarrer son réseau IP	109
Le plan d'adressage IP	110
La démarche	110
Les principes de base	112
Impact sur l'Internet	113
Les sous-réseaux IP	113
Méthode d'affectation des réseaux LAN	115
Méthode d'affectation des réseaux WAN	118
Méthode d'affectation des stations au sein des réseaux	119
L'encapsulation des protocoles	120
L'adressage	121
Le multiplexage	122
Comment une station envoie-t-elle une trame Ethernet à une autre ?	124
Échange de trames sur un segment Ethernet	125
Échange de trames entre différents segments Ethernet	126
Comment une station envoie-t-elle un paquet IP à une autre ?	131
La résolution d'adresse	134
Comment une application envoie-t-elle des données ?	135

CHAPITRE 8

Mettre en place sa première interconnexion de réseaux	139
Le contexte	140
Les choix de base	140
Quel support de transmission ?	140
Quel protocole de niveau 2 ?	141
Quel équipement réseau ?	142
Quel opérateur ?	142
De quoi avons-nous besoin ?	143
D'une liaison entre les deux sites	143

<i>À quel débit ?</i>	143
<i>Avec quel support de transmission ?</i>	144
<i>Avec quel service opérateur ?</i>	145
De routeurs	146
De câbles	147
Comment faire fonctionner tout cela ?	148
Définir l'architecture	148
Connecter un PC au routeur	149
Configurer le routeur	150
<i>Affecter les adresses IP</i>	150
<i>Activer le routage</i>	151
Configurer les postes de travail	153
Tester le réseau	156
Optimiser	156
Mettre en place une liaison de secours	157
Quels sont les choix ?	157
Solutions alternatives	157
Installation d'un accès de base T0	160
Sécurisation de la liaison	161
Gestion du débordement	162

CHAPITRE 9

Architecture des réseaux étendus	163
Les solutions disponibles sur le marché	164
Les infrastructures	164
Les réseaux opérateurs	165
L'accès au réseau	165
Les services proposés par les opérateurs	165
Les choix du client	167
Le réseau de transport	169
Qu'est-ce qu'une LS ?	169
La boucle locale	171
Les applications xDSL	175

Dimensionner les liaisons	176
Identifier les flux	176
<i>Les flux de type conversationnel</i>	178
<i>Les flux de type transactionnel</i>	178
<i>Les flux de type transfert de fichiers</i>	179
<i>Les flux client-serveur</i>	180
Estimer la volumétrie	180
<i>Volumétrie liée à la messagerie</i>	181
<i>Volumétrie liée aux transferts de fichiers</i>	182
<i>Volumétrie liée aux applications transactionnelles site central</i>	182
<i>Volumétrie liée aux applications transactionnelles web</i>	182
<i>Volumétrie liée à d'autres services</i>	182
<i>Rassembler toutes les données</i>	183
Calculer les débits	184
<i>Tenir compte des temps de réponse</i>	185
CHAPITRE 10	
Bâtir un réseau de transport	187
LS, Frame Relay ou ATM ?	188
Mettre en place un réseau de LS	189
Mettre en place un réseau Frame Relay	190
Qualité de service et facturation	192
<i>Débit garanti</i>	192
Connecter un routeur au réseau de transport	194
<i>Si le routeur ne supporte pas Frame Relay</i>	196
<i>Si le routeur supporte Frame Relay</i>	197
Gérer les circuits virtuels	198
Combien de circuits virtuels ?	199
Configurer les PVC	201
<i>Correspondance entre adresses IP et DLCI</i>	202
Configurer les SVC	203
Gérer la qualité de service	204
Les sous-interfaces	206
Mettre en place un réseau ATM	207

Qualité de service et facturation	208
<i>La gestion du trafic : TMS 4.0 (ATM Forum af-tm-0056.000)</i>	208
<i>Les classes de service ATM Transfer Capabilities (ITU I.371 et TMS)</i>	209
Connexion du routeur au réseau de transport	210
<i>Si le routeur ne dispose pas d'interface ATM</i>	211
<i>Si le routeur supporte ATM</i>	212
Configurer les SVC	214
Gestion de la qualité de service	218
Les paramètres décrivant les classes de service (ITU I.356 et ATM Forum TMS 4.0)	219
L'adressage	220
L'adressage NSAP (ISO 8348, ITU X.213, RFC 1629)	220
L'adressage ATM	221
L'adressage Frame Relay	222
Interopérabilité entre Frame Relay et ATM	222

CHAPITRE 11

Assembler les briques du LAN et du WAN	223
Mise en place un réseau fédérateur	224
Les données du problème	224
La démarche	224
Quelle technologie ?	224
Quels équipements ?	225
Routeur ou commutateur de niveau 3 ?	226
Quelle architecture ?	227
Configurer les VLAN	228
Extension du réseau fédérateur	231
L'adressage et le routage IP	233
Redondance du routage	233
La rencontre du LAN et du WAN	236
Le routage sur le WAN	237
Configuration du routage	237
Redondance en cas de panne	239
Ajustement des paramètres	241

<i>Diffuser les routes statiques</i>	241
<i>Modifier le coût des routes</i>	241
<i>Limiter la diffusion des routes</i>	242
<i>Modifier la fréquence des échanges</i>	242
<i>Forcer l'élection du routeur désigné</i>	243
Les performances d'OSPF	243

TROISIÈME PARTIE SE PRÉPARER AU MULTIMÉDIA

CHAPITRE 12

Les flux multimédias	247
Les caractéristiques des flux multimédias	248
Les codec audio	249
Les codec vidéo	250
Les problèmes posés par les transmissions audio et vidéo	252
Estimation du temps de transit	253
Le transport des données multimédias	254

CHAPITRE 13

Le routage des flux multimédias	257
La diffusion sur un réseau IP	258
La gestion des groupes de diffusion	260
Le routage des flux multicast	264
Le routage à l'aide de DVMRP	264
Le routage à l'aide de MOSPF	269
Le routage à l'aide de PIM	274
<i>Principe de PIM-SM</i>	275
<i>Principe du calcul des routes</i>	275
<i>Principe du routage</i>	277
<i>Routage sur les liaisons WAN</i>	277
Quel protocole choisir ?	279
Architecture adaptée au protocole DVMRP	281

Architecture adaptée au protocole MOSPF	281
Architecture adaptée au protocole PIM	282
Contrôler la diffusion sur son réseau	283

CHAPITRE 14

La qualité de service sur IP	287
Améliorer les performances du réseau	288
Affecter des priorités sur les files d'attente	288
Agir sur les files d'attente	290
<i>L'algorithme FIFO - Un fonctionnement simple</i>	290
<i>Gérer les congestions</i>	290
<i>Prévenir les congestions</i>	291
<i>Réguler le trafic</i>	291
Quelle file d'attente choisir pour son réseau ?	293
Gérer la qualité de service	294
La qualité de service selon DiffServ	294
Le champ TOS	294
Configuration des routeurs	296
Configuration des commutateurs de niveau 2	298
Configuration des commutateurs de niveau 3	300
<i>Définir une règle de marquage</i>	300
<i>Définir une règle de policing</i>	300
<i>Définir une règle de classification</i>	301
<i>Associer une politique à un port</i>	302
<i>Affecter des valeurs au champ DSCP</i>	302
Configuration des postes de travail	303
La qualité de service selon IntServ	305
La réservation des ressources	307
La description de la qualité de service	312
<i>Les classes de service</i>	312
<i>Description des classes de service</i>	313
<i>Caractéristique des flux</i>	313
<i>Objets RSVP</i>	313

Définir une politique de qualité de service	315
CHAPITRE 15	
La téléphonie et la vidéo sur IP	317
Présentation des protocoles multimédias	318
Les composants d'un système H.323	319
L'établissement d'une communication	322
Interconnecter les PABX <i>via</i> IP	324
Mettre en place un gatekeeper	328
La voie vers le tout IP	333
Configurer le PABX et la passerelle VoIP	335
Déclarer les terminaux téléphoniques	336
Assurer la qualité de service	338
Transporter les flux multimédias	339
Le transport des flux audio et vidéo via RTP et RTCP	340
Optimiser les flux multimédias	343
Compression des en-têtes	343
Utilisation des mixers	344
Échanger des données multimédias	345
QUATRIÈME PARTIE GÉRER SON RÉSEAU	
CHAPITRE 16	
Administrer son réseau IP	349
Les utilitaires de base	350
Le ping	350
Le traceroute	351
Observer ce qu'il se passe sur son réseau	354
Piloter son réseau	356
Quelle station d'administration ?	356
Pour quelle utilisation ?	357

Configurer automatiquement ses PC	363
Quelle utilisation de DHCP ?	363
Comment configurer un serveur DHCP ?	365
<i>Définir les pools d'adresses</i>	365
<i>Définir les options à distribuer</i>	367
Configurer les routeurs	370
Installer plusieurs serveurs	371
Vérifier la configuration de son PC	371

CHAPITRE 17

La gestion des noms	375
Présentation du DNS	376
Les composants du DNS	376
Élaborer un plan de nommage	376
Définir l'arborescence DNS	377
Standardiser le nommage des objets	379
Configurer les serveurs DNS	381
Configurer le fichier cache	384
Configurer un serveur primaire	385
<i>Activer la résolution de nom</i>	388
<i>Activer le routage de la messagerie</i>	389
<i>Du bon usage des alias</i>	390
Configurer un serveur racine	390
Configurer un serveur secondaire	392
Configurer un serveur cache	392
Déléguer l'autorité à un autre serveur	393
Les domaines de résolution inverse	393
Le fichier d'initialisation	394
Configurer les clients DNS	395
Vérifier le fonctionnement du DNS	396

Annexes

Normes et standards	401
Le câblage	401
<i>Normes CEN relatives au câblage</i>	401
<i>Normes EIA/TIA relatives au câblage</i>	401
<i>Normes ITU-T relatives au câblage</i>	401
Les interfaces physiques	402
<i>Avis de l'ITU-T relatifs aux interfaces physiques</i>	402
<i>Normes EIA/TIA relatives aux interfaces physiques</i>	402
<i>Avis de l'ITU-T relatifs aux échanges ETDD-ETCD</i>	402
Les réseaux locaux	403
<i>Normes IEEE relatives aux réseaux locaux</i>	403
La famille des protocoles TCP/IP	404
<i>RFC relatives aux protocoles TCP/IP</i>	404
<i>Standards originaux du DOD (Department Of Defense) relatifs à TCP/IP</i>	405
<i>RFC relatives aux protocoles de routage IP</i>	405
<i>RFC relatives aux applications utilisant TCP/IP</i>	405
<i>RFC relatives à IP sur Frame-Relay</i>	406
<i>RFC relatives à IP sur ATM</i>	406
<i>RFC relatives à PPP</i>	407
<i>RFC relatives à SNMP</i>	407
<i>Normes ISO et équivalents ITU-T relatifs à la syntaxe ASN.1</i>	408
<i>RFC relatives à IPv6</i>	408
Le multimédia sur IP (VoIP)	409
<i>RFC relatives à la voix sur IP</i>	409
<i>Avis de l'ITU-T relatifs à la voix sur IP</i>	409
<i>RFC relatives à la qualité de service</i>	410
<i>RFC relatives au routage multicas</i>	410
Les réseaux RNIS	411
<i>Organisation et nomenclature des normes relatives</i> <i>aux réseaux numériques à intégration de services</i>	411
<i>Série I.100 : Concepts généraux du RNIS</i>	412
<i>Série I.200 : Services assurés par le RNIS</i>	412
<i>Série I.300 : Aspects réseaux du RNIS</i>	412

<i>Série I.400 - Interfaces usager-réseau</i>	413
<i>Série I.500 : Interfaces d'interconnexion du RNIS</i>	414
<i>Série I.600 : Administration du RNIS</i>	414
<i>Avis de l'ITU-T relatifs aux réseaux ATM</i>	415
<i>Avis de l'ITU-T et équivalents ANSI relatifs au relais de trames</i>	415
<i>Avis de l'ITU-T relatifs aux systèmes de transmission numérique MIC</i>	416
<i>Avis de l'ITU-T relatifs aux réseaux SDH</i>	416
 Glossaire	 417
 Bibliographie	 433
 Sites web	 435
Câblage	435
Internet	435
Modem-câble	436
Organismes de normalisation	436
Protocoles	436
Qualité de service	437
Réseaux sans fils	437
Revue de presse	437
VoIP	438
 Index	 439
 Table des encarts	 445

1

Introduction

J'ai retrouvé, il y a peu, un ouvrage faisant partie de ceux qui ont abreuvé une génération d'étudiants en réseaux & télécom, un pavé de plus de 900 pages. Il faisait partie d'une de ces bibles immanquables que l'on se devait de lire. Celui qui n'avait pas lu le « Machin » ou le « Truc » passait assurément à côté de quelque chose, et risquait de compromettre ses examens.

Sur les 900 pages dédiées aux réseaux, une seule était consacrée à TCP/IP sous une rubrique bizarrement intitulée « La productique ». Il y était dit que cette architecture était démodée et que, depuis quelques années, tous les utilisateurs de ce type de réseau étaient amenés à évoluer vers l'architecture OSI. Le livre datait de 1987, TCP/IP avait 18 ans et l'Internet explosait... aux États-Unis.

Cela m'a rappelé les quelques temps passés à travailler au cœur de la Silicon Valley en tant que programmeur. J'étais alors en charge de développer des couches logicielles autour de TCP/IP. Un de mes collègues avait affiché à l'entrée de son bureau un manifeste intitulé « Why OSI ? ». Parmi les réponses saugrenues, il y avait celles-ci : « parce que c'est normalisé », « parce qu'il y a 7 couches », « parce que c'est compliqué », etc.

Voilà un des problèmes de l'Europe : d'un côté un centre d'activité qui crée la technologie de demain, de l'autre des commentateurs avertis. Dans l'entreprise, contentons-nous donc d'utiliser au mieux ce qu'on nous propose.

J.-L. Montagnier

Le principe de l'apport minimal

L'histoire montre que les technologies qui se sont imposées sont celles qui répondent à au moins un des trois critères suivants : conservation de l'existant, réponse aux besoins et réduction des coûts.

Par exemple :

La nouvelle technologie	
Conserve l'existant ?	Oui à 50 %
Répond aux besoins ?	Oui à 50 %
Réduit les coûts ?	Oui à 10 %
TOTAL	+ de 100 %

→ Oui, ça apporte quelque chose !

La nouvelle technologie peut ne pas conserver l'existant, mais apporter une réelle plus-value ; c'est le cas du Compact Disc qui a remplacé le vinyle en quelques années (meilleure qualité de son, plus grande résistance , etc.).

La nouvelle technologie peut être plus chère mais répondre aux besoins ; c'est le cas des téléphones mobiles. Pour un usage personnel, on a besoin de communiquer, de se sentir important (besoins irrationnels). Pour un professionnel, être joignable à tout moment fait partie de la qualité de service qu'il offre à ses clients (besoin rationnel) : cela lui permet d'être plus réactif et donc de gagner plus d'argent, même si les communications coûtent deux fois plus cher que celles d'un téléphone fixe.

La nouvelle technologie peut remettre en cause l'existant mais réduire les coûts. Si le retour sur investissement est assuré, elle a alors de fortes chances de s'imposer. Ce constant souci de productivité est présent dans toutes les industries.

Face à une nouvelle technologie, la question à se poser est donc : est-ce qu'elle apporte quelque chose ?

Un exemple dans le domaine des réseaux

L'Ethernet à 100 Mbit/s, puis le Gigabit Ethernet se sont imposés très rapidement, parce que d'une part les composants électroniques proviennent de technologies existantes (*Fibre Channel*, etc.) ce qui réduit les coûts, et d'autre part l'existant est préservé.

ATM est sans doute ce qui se fait de mieux en matière de technologie réseau, mais il en fait trop par rapport aux besoins d'aujourd'hui. En faire trop implique de dépenser plus d'argent en R&D, en formation, en fabrication, etc., ce qui a pour conséquence de ralentir la diffusion de ladite technologie.

Les partisans des deux technologies placent le débat au niveau de la qualité de service, de l'intégration voix/données/image et des débits.

A-t-on besoin d'une qualité de service ?

En termes de gestion de la qualité de service, ATM est équivalent à Frame Relay ou IP puis-
sance 10. Mais, actuellement, 10 % des capacités techniques d'ATM sont exploitées : la
classe de service la plus utilisée est l'ABR (qui est la plus basse disponible !), l'accès le plus
utilisé est AAL-5 (qui offre le minimum de service !), et l'intégration voix/données se fait
attendre.

La gestion de la qualité de service devient nécessaire lorsque l'on doit gérer l'utilisation
d'une bande passante limitée. Or, toutes les évolutions actuelles montrent que les débits ré-
seaux ne cessent d'augmenter et que la bande passante devient disponible à volonté et à
moindre coût.

Dès lors que la bande passante est disponible à moindre frais, une gestion évoluée de la QoS
devient inutile. Les entreprises ont juste besoin d'une QoS de base : **priorité, débit garanti,
délai de transit et variation du délai de transit** (la question de la perte de données ne se
pose même plus avec des taux d'erreur 10^{-6} à 10^{-10} sur les lignes numériques).

Par contre, seuls les opérateurs y trouvent leur compte : ATM leur permet de réduire leurs
coûts de fonctionnement en dimensionnant leur réseau au plus juste tout en y casant le
maximum de clients. Une gestion évoluée de la QoS leur permet également de présenter des
grilles tarifaires complexes.

L'intégration voix/données/image

Quelles sont les technologies que les entreprises utilisent aujourd'hui pour la voix, les don-
nées et l'image ? Réponse : Frame Relay et IP. Bien sûr, en prenant en compte l'existant, on
trouve encore du X.25, du FDDI, du Decnet, du SNA, etc.

Jusque vers 1996, le monde de la téléphonie et les partisans d'ATM (souvent les mêmes)
doutaient que l'on puisse transporter de la voix sur Frame Relay et dans des paquets IP. Au-
jourd'hui, on ne peut que constater la réalité.

Pourquoi ? Parce que les technologies ont évolué ; elles permettent de faire plus de choses à
moindre coût et répondent toujours au principe de l'apport minimal.

Aujourd'hui, IP permet de transporter de la voix compressée d'assez bonne qualité en utili-
sant donc moins de débit. Les systèmes de visioconférence utilisent majoritairement le
RNIS, Frame Relay et, dans une moindre mesure, IP.

Les débits

Jusqu'à la fin des années 90, ATM était annoncé comme étant le seul à pouvoir offrir des
débits élevés, de 155 à 622 Mbit/s, voire plus dans le futur.

Jusqu'à ce que le Gigabit Ethernet arrive. Dommage pour ceux qui ont investi dans ATM à 25 Mbit/s.

Aujourd'hui, ATM est utilisé par les entreprises dans quelques rares cas :

- dans les réseaux fédérateurs et dans les réseaux de campus ;
- dans les mondes médical et audiovisuel pour la vidéo.

Le Gigabit Ethernet s'est par contre imposé rapidement dans les réseaux fédérateurs. En effet, outre son meilleur rapport qualité/prix, il est compatible avec l'existant tout en simplifiant les architectures.

Le cas ATM

La technologie ATM (*Asynchronous Transfer Mode*) a toujours été présentée comme un aboutissement, comme seule capable d'assurer le transport des données multimédias. Elle n'a cependant jamais réussi à s'imposer dans les entreprises et chez les particuliers. Seuls les opérateurs en font un usage important. Alors ?

Pourquoi ATM pourrait s'imposer ?

Parce que les opérateurs et les constructeurs investissent massivement dans cette technologie. Mais est-ce que cela assurera sa diffusion dans les entreprises et chez les particuliers ?

Parce que, avec 1 Gbit/s, Ethernet arrive au bout de ses limites. Mais est-ce vrai ? Alors qu'ATM plafonne encore à 622 Mbit/s, les ingénieurs se penchent aujourd'hui sur l'Ethernet à 10 Gbit/s !

Parce qu'ATM fonctionne aussi bien sur LAN que sur WAN. Mais le Gigabit aussi !

Parce qu'ATM intègre la voix, la vidéo et les données. Mais qui utilise aujourd'hui cette possibilité ?

Parce qu'ATM s'intègre dans l'existant avec *LAN Emulation*. Mais que c'est compliqué !

Pourquoi ATM pourrait disparaître ?

Parce que son coût n'arriverait pas à baisser suffisamment. Entre une carte Ethernet 100 Mbit/s à 700 F et une carte ATM 155 Mbit/s à 3 000 F, on n'hésite pas.

Parce qu'ATM engendre un *overhead* d'au moins 10 % (5 octets d'en-tête pour 48 octets de données dans le meilleur des cas avec AAL-5). Même si la bande passante est disponible à volonté, il y a un surcoût de 10 %, ce qui représente un manque de compétitivité important pour une entreprise. Les apports fonctionnels (QOS, intégration voix/données, etc.) justifient-ils cet *overhead* ? Le compensent-ils ?

Enfin, il est possible d'établir des liaisons Ethernet Gigabit sur de longues distances (actuellement, 150 km). C'est l'occasion d'utiliser l'Ethernet de bout en bout (simplification de l'architecture).

ATM peut-il remplacer IP ?

Cette question souvent brandie comme une menace par la partisans d'ATM sème la confusion dans les esprits, car IP et ATM ne remplissent pas le même rôle au sein d'un réseau, même si certaines de leurs fonctions se recouvrent.

ATM offre, en effet, les mêmes fonctions qu'IP : une signalisation, un adressage (NSAP), un routage (PNNI) et la gestion de la qualité de service.

Mais le transport des données entre un PC et un serveur s'effectue exclusivement à l'aide du couple TCP/IP. Qui transfère des fichiers sur ATM sans passer par IP ? En fait, ATM seul ne sert pas à grand-chose.

Pis, le protocole IP peut se passer des services d'ATM : le plus souvent, IP circule sur PPP, voire directement sur la fibre optique (en SDH ou encore plus directement en WDM).

Pourquoi alors remplacer une technologie par une autre ?

Le syndrome de Vinci

Faut-il alors donner du temps au temps ?

Par exemple, IP n'a pas connu une diffusion planétaire immédiate : ce protocole s'est imposé rapidement dans les universités et les centres de recherche, puis plus lentement dans le monde des entreprises.

ATM pourrait suivre la même voie : il fait l'objet de nombreux travaux de normalisation et a le soutien de l'ONU (l'ITU) et des constructeurs (ATM Forum).

Cela n'est cependant pas un gage de réussite ; il suffit de se souvenir de ce qui est arrivé aux protocoles OSI et au RNIS. Mais certains peuvent encore espérer que le RNIS sera le Concorde qui aura permis le développement de l'Airbus ATM.

Le problème des technologies trop en avance sur leur temps est qu'elles sont soumises au syndrome de Vinci¹ :

- Au début, elles ne répondent pas au principe de l'apport minimal.
- Dix ans après, elle sont dépassées (techniquement ou économiquement) par d'autres technologies ou de plus anciennes qui ont évolué avec leur temps.

C'est ce qui aurait pu arriver à IP si les protocoles OSI avait été plus performants, et c'est ce qui pourrait arriver à ATM dans quelques années.

À court et moyen termes, ATM ne s'imposera pas face à IP et à Ethernet. Dans dix ans, on pourra se poser de nouveau la question : Ethernet et IP ont-ils atteint leurs limites par rapport

¹ Artiste de génie, Léonard de Vinci était également « Premier ingénieur et architecte du Roi, Mécanicien d'état ». Cependant, nombre de ses projets sombraient dans les limbes, parce que trop novateurs ou irréalisables avec les moyens techniques de l'époque. Ses talents d'ingénieur étaient surtout mis à contribution pour réaliser des automates de fêtes foraines et de spectacles.

aux besoins du moment ? Si oui, quelle est la solution qui répond au principe de l'apport minimal ?

On est par ailleurs étonné du parti pris de la presse pour l'ATM. Les revues spécialisées ont régulièrement émis des avis négatifs à propos des commutateurs 100bT, puis du Gigabit :

- juillet 1997 : « qualité de service à revoir » ;
- janvier 1998 : « fonctions d'exploitation limitées » ;
- etc.

Ces mêmes magazines mettent en avant le handicap de la jeunesse du Gigabit ou encore sa limitation à 1 Gbit/s (en oubliant que l'on attend toujours l'ATM à 622 Mbit/s pour les réseaux locaux...). D'une manière générale, on met en exergue les avantages de l'ATM par rapport au Gigabit en oubliant les inconvénients du premier et les avantages du second.

En réalité lorsqu'ils bénéficient d'une dynamique importante comme Ethernet et IP, les technologies et les produits s'améliorent au fil du temps. La méthode américaine est, en effet, de sortir un produit le plus rapidement possible afin rentabiliser les premiers investissements, les améliorations ne venant que si le marché se développe.

Les technologies gagnantes

Quelles sont elles ? Celles qui sont de plus en plus utilisées dans les entreprises. Celles qui ont supplanté les autres.

Le monde des réseaux a, en effet, longtemps été caractérisé par une quantité innombrable de protocoles plus ou moins exotiques, essentiellement poussés par des constructeurs soucieux de verrouiller leur marché (Digital avec Decnet, IBM avec SNA, Novell avec IPX, etc.) et par des organismes de normalisation sacrifiant aux plaisirs des techniciens (OSI de l'ISO...).

Ainsi, seuls quelques protocoles ont survécu ; ils ont pour point commun d'avoir su s'adapter aux besoins des entreprises tout en présentant le meilleur rapport qualité/prix.

Par exemple, l'Ethernet de l'an 2000 ne ressemble plus à celui des années 70. La famille TCP/IP s'est enrichie de dizaines de protocoles. Et le Frame-Relay a su intégrer la voix et les données en offrant le minimum de qualité de service nécessaire. Tandis que le RNIS a su répondre à un besoin bien spécifique d'interconnexion.

Tout cela parce ces protocoles reposent sur les technologies ouvertes, simples et qui apportent réellement une plus-value aux entreprises.

Cet ouvrage s'attache donc à décrire ces protocoles gagnants que sont **Ethernet**, **TCP/IP** et **Frame-Relay**.

Il décrit également l'utilisation spécifique que les entreprises font du **RNIS** (pour l'interconnexion des LAN) et d'**ATM** (pour son usage dans le WAN et également pour vous montrer combien c'est compliqué...).

Il s'attache également à décrire des protocoles émergents : **H.323** dans le domaine de la voix et de l'image sur IP, **IGMP** et **PIM** pour le routage multicast, **RSVP**, **IntServ** et **Diff-Serv** pour la qualité de service, ainsi que les technologies **xDSL** pour l'accès aux réseaux des opérateurs et à l'Internet.

D'autres sont à venir et pourraient trouver leur place dans les entreprises (et dans cet ouvrage !) : **SIP** (concurrent de H.323), **GMRP** (multicast sur les commutateurs Ethernet), **MPLS** (commutation de paquets IP) ou encore **IPv6** (IP nouvelle génération) et **SNMPv3** (administration des réseaux).

Cet ouvrage présente également quelques systèmes bien répandus dans les entreprises, tels que le navigateur web (protocole **HTTP**), le transfert de fichiers **FTP**, le **DHCP** (configuration des PC) et le **DNS** (service de noms). D'autres utilitaires et protocoles annexes sont également décrits.

Enfin, le **câblage**, sur lequel tout repose, n'est pas oublié.

De nombreux protocoles, qui existent pourtant toujours, ne sont donc pas traités : **ATM** (pour son usage dans le LAN), **FDDI**, **Token-Ring**, **X.25**, **SNA** (et les protocoles associés, tels que **DLSW** et **STUN**).

Le **Token-Ring** a, en effet, été supplanté par l'**Ethernet**, le **X.25** par le **Frame-Relay**, **IPX** par **IP**, etc. Le **SNA** n'est présent que dans les sociétés qui ont beaucoup investi dans les *main-frame* IBM. La migration vers **IP** est cependant bien amorcée.

Les évolutions

Il est toujours risqué de se livrer à des pronostics. Mais allons-y quand même.

Ethernet et **TCP/IP** focalisent plus que jamais les efforts en recherche et en développement. Ils continueront donc d'évoluer. Le **Frame-Relay**, longtemps considéré comme étant une technologie de transition vers l'**ATM**, ne perdurera que s'il peut s'adapter aux besoins futurs.

À force d'investir, les constructeurs et opérateurs parviendront peut-être à imposer **ATM** sur le **WAN**. Par exemple, l'accès Internet *via* **ADSL** que France Télécom propose aux particuliers utilise **ATM**.

L'évolution de fond concerne la voix sur **IP** : consommant moins de bande passante, permettant d'utiliser le réseau **IP** existant, elle est source d'économies importantes. Bien que ne répondant pas à un réel besoin (le téléphone, aujourd'hui, ça marche), elle apporte l'image. Après le téléphone portable, le visiophone pourrait être le prochain gadget à la mode.

Le passage au tout **IP** dans les entreprises ne sera possible que si les nouveaux **PABX** sont plus ouverts et plus accessibles que leurs équivalents traditionnels. Un premier frein à cette extension pourrait venir des constructeurs traditionnels qui veulent conserver leur pré carré. Bien qu'ouverte à **H.323**, leur offre conserve bien des aspects propriétaires. Les constructeurs informatiques, Cisco en tête, pourraient bien en profiter.

Une autre évolution importante semble se dessiner : la distinction traditionnelle entre LAN (réseau local), MAN (réseau métropolitain) et WAN (réseau longues distances) est de moins en moins vraie. Aujourd'hui, on trouve Ethernet dans ces trois parties du réseau. Certains opérateurs proposent ainsi des liaisons Gigabit Ethernet sur plus de 150 km. Le Gigabit pourrait donc détrôner l'ATM sur ce segment.

Enfin, les technologies xDSL ratissent large : elles regroupent, en effet, différents protocoles adaptés à différents besoins. L'ADSL, la version économique, pourrait s'imposer chez les particuliers. C'est une question de coût. L'HDSL, la version haut de gamme pour les entreprises, pourrait être imposée par les opérateurs.

PREMIÈRE PARTIE

Se lancer dans les réseaux

Premier contact avec les réseaux : l'Internet

Qu'est-ce qu'un réseau ? Et, plus exactement, un réseau informatique, un réseau de télécommunications ? Pour l'expliquer de manière pratique, il m'a paru plus simple de débiter par un réseau dont tout le monde (ou presque) a entendu parler : l'Internet, le plus grand réseau public, le réseau des réseaux. Qui dit réseau dit connexion, ce qui implique l'utilisation d'un ordinateur ou de tout autre terminal, tel un téléphone portable.

Ce chapitre décrit votre premier contact avec le monde des réseaux, à partir de cet instant où vous vous connecterez à un serveur depuis votre domicile ou votre entreprise.

Vous y apprendrez notamment :

- les principes de base des réseaux ;
- ce qu'est un modem, un câble série ;
- ce qu'est un driver, une couche réseau ;
- comment paramétrer les logiciels de communication.

Le but du jeu est le suivant : vous possédez un ordinateur (de préférence un PC, cela représente quand même près de 90 % du marché) et vous voulez vous connecter à l'Internet.

– Pour quoi faire ?

– Eh bien...

– Eh bien, pour dialoguer avec le reste du monde, chercher des informations de toute sorte, bref, s'amuser et apprendre.



Qu'est-ce que l'Internet ?

L'Internet est, avant tout, une collection de millions de réseaux auxquels sont connectés des millions d'ordinateurs. Personne n'est propriétaire de l'Internet dans son ensemble.

Le cœur de l'Internet est constitué par la concaténation des réseaux des opérateurs de télécommunications (France Télécom, AT&T, MCI-Worldcom, etc.).

Un réseau est donc caractérisé par un aspect **physique** (les câbles véhiculant des signaux électriques) et un aspect **logique** (les logiciels qui réalisent les protocoles).

À la périphérie gravitent les fournisseurs d'accès, également appelés les ISP (*Internet Service Provider*), tels qu'Oléane, AOL (*American On Line*), Compuserve, Club-Internet, etc. Ce sont les ISP qui permettent aux sociétés et aux particuliers de se connecter à l'Internet *via* une liaison permanente ou une liaison téléphonique. Aujourd'hui, quasiment tous les ISP sont affiliés à des opérateurs.

Les opérateurs disposent de réseaux plus ou moins identiques : il se peut que les liaisons passent par les mêmes câbles ou empruntent des satellites ou encore que ces réseaux soient interconnectés à plusieurs endroits. D'ailleurs, les opérateurs louent des portions de leur réseau à d'autres opérateurs, qui peuvent ensuite les sous-louer à d'autres plus petits.

Enfin, l'Internet est également constitué de grands réseaux issus d'organismes publics (centres de recherches, universités, gouvernement, etc.) tels que Renater (réseau de la recherche française), Abilene (réseau de la recherche américaine), etc.

QU'EST-CE QU'UN RÉSEAU ?

D'une manière générale, un réseau est un système permettant de relier des ordinateurs entre eux.

Il est constitué d'un ensemble de **câbles** en cuivre et en fibre optique véhiculant des signaux. Un signal représente une unité d'information (le **bit**) émise par un ordinateur. Une série de bits permet d'identifier de manière unique une information, par exemple la lettre "A".

Comme pour les humains, il est nécessaire d'établir des règles de communication du genre : – Bonjour, comment ça va ? – Bien, et toi ? – Bien. As-tu reçu mon message ? – Non. – Ah bon ? Attends, je te le renvoie.

C'est ce qu'on appelle un **protocole** de communication. Pour cela, les ordinateurs disposent de logiciels spécialisés.

Un réseau est donc caractérisé par un aspect **physique** (les câbles véhiculant des signaux électriques) et un aspect **logique** (les logiciels qui réalisent les protocoles).

De quoi a-t-on besoin ?

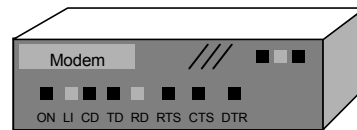
D'un modem

Le premier élément est le modem (**mod**ulateur-**dé**modulateur). Cet appareil permet de convertir les informations numériques émises par l'ordinateur en signaux transmissibles sur de longues distances.

On distingue plusieurs catégories de modems selon le support de transmission utilisé. Celui dont il s'agit ici est de type analogique ; il permet d'utiliser le réseau téléphonique classique.

Le modem de base se présente sous la forme d'un petit boîtier, comportant généralement deux ou trois connecteurs et des voyants rouges et/ou verts.

Figure 2-1.



LED indiquant l'état du modem :

LI	Line
CD	Carriage Detect
TD	Transmit Data
RD	Receive Data
RTS	Request To Send
CTS	Clear To Send
DTR	Data Transmit Ready

QU'EST-CE QU'UN MODEM ?

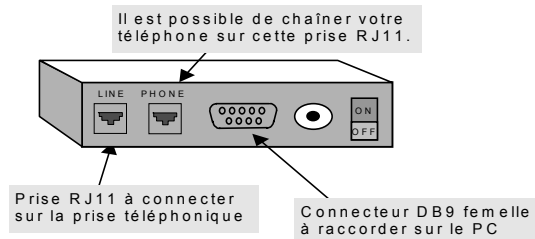
L'ordinateur traite les informations sous forme **numérique** et sous forme de mots de 32 ou 64 bits (64 éléments d'informations binaires : 0 et 1).

Il est capable d'émettre directement ces mots à haute vitesse, mais sur de courtes distances (quelques mètres). Les modems permettent d'émettre sur de longues distances, mais à des débits plus faibles, en utilisant des techniques de **modulation du signal** : la valeur d'un bit (0 ou 1) prend la forme d'une onde électrique, et le signal obtenu est alors de type **analogique**.

Le réseau téléphonique transporte la voix sur deux fils, ce qui permet de ne transmettre qu'un bit à la fois. Les mots de 32 et 64 bits émis par un ordinateur sont donc transmis en **série** et de manière **asynchrone** (l'émetteur et le récepteur ont des horloges différentes).

La vitesse de transmission est exprimée en bauds, ou en kilo-bits par secondes (**Kbit/s**).

Les modèles actuels offrent un débit maximal de transmission qui est de 56 Kbit/s (norme V.90). En général, un modem de ce type supporte toutes les normes inférieures (et historiquement les premières) jusqu'au V.23 du Minitel (1 200 bit/s dans un sens et 75 bit/s dans l'autre).

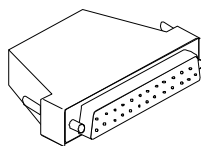


D'un câble série

Pour connecter le modem, l'ordinateur doit disposer d'une prise appelée indifféremment connecteur, port série ou port asynchrone. Ce connecteur est de type DB9 ou DB25 femelle côté modem, et DB9 mâle côté PC. On trouve également des connecteurs de type DB25 (le numéro indique le nombre de broches — *pin*, en anglais — sur le connecteur).

Figure 2-2.

Connecteurs utilisés pour les liaisons série.



Connecteur DB25 femelle



Connecteur DB9 mâle

Par extension, le câble servant à relier l'ordinateur et le modem est appelé câble série. On l'appelle également câble croisé, car il présente la particularité suivante : le fil qui part de la broche "émission" d'un côté est raccordé à la broche "réception" de l'autre, et inversement. Dans la pratique, les fils "émission" et "réception" se croisent au niveau de l'un des deux connecteurs.

D'une prise téléphonique

Le moyen le plus simple de se connecter à l'Internet est de passer par le réseau téléphonique commuté, connu sous le sigle RTC. Dédies à l'origine pour véhiculer de la voix, les modems permettent de moduler des signaux pour transmettre des données. Lorsque le modem établit la communication et que le haut-parleur est activé, vous entendez un court instant le chuintement caractéristique qui en résulte. La transmission de la parole ne procède pas autrement : le téléphone sert à convertir les ondes sonores en signaux électriques.

Le RTC est donc le support de transmission le plus simple pour se connecter à l'Internet. D'autres supports seront étudiés dans ce livre : modem câble, modem ADSL, modem RNIS et lignes spécialisées.

Notez que vous pouvez utiliser votre prise téléphonique pour y raccorder à la fois le téléphone et le modem, mais que vous ne pouvez pas les utiliser en même temps.

QU'EST-CE QUE LE RTC ?

Le RTC (Réseau téléphonique commuté) est tout simplement le réseau auquel est connecté votre téléphone.

Il est constitué d'autocommutateurs (**autocom** en abrégé) présents sur toute la France, reliés entre eux par des liaisons à haut débit véhiculant simultanément des milliers de conversations téléphoniques.

Votre téléphone est directement connecté à un autocom via une paire de fils en cuivre. Cette partie du réseau est indifféremment appelée desserte locale, desserte de raccordement ou encore **boucle locale**.

Le réseau est dit commuté, ou à **commutation de circuits**, car l'établissement d'une communication consiste à activer des circuits électroniques (autrefois électromécaniques) au sein des autocoms.

D'un abonnement à l'Internet

Vous êtes chez vous. Mais où est l'Internet ? Réponse : partout dans le monde et nulle part en particulier. Le réseau Internet passe sans doute devant chez vous, dans une fibre optique enterrée, mais, malheureusement, vous ne pouvez pas vous y connecter directement (se serait vite l'anarchie). Il faut passer par l'intermédiaire d'un fournisseur d'accès à l'Internet, un des fameux ISP (*Internet Service Provider*).

Ces derniers disposent de ce que l'on appelle des points de présence ou POP (*Point Of Presence*). Un POP est une salle informatique équipée d'ordinateurs pilotant des pools de modems raccordés au RTC. Ces ordinateurs sont eux-mêmes connectés à l'Internet. Nous y voilà enfin !

L'abonnement en lui-même coûte entre 0 et 90 francs par mois selon le forfait retenu. Attention, pour tous les abonnements gratuits le coût des communications téléphoniques est à votre charge (il est le même que lorsque vous téléphonez). Pour diminuer la facture, il est donc impératif de se connecter au POP le plus proche de chez vous. Le mieux est d'en choisir un qui est situé à moins de 10 km, afin de bénéficier du coût d'une communication locale (le moins cher). Le principal critère de choix d'un ISP est donc la proximité des POP. Dans la pratique, votre ISP met à votre disposition un numéro unique (de type 0860), valable pour toute la France, qui permet de router automatiquement votre appel vers le POP le plus proche.

Le deuxième critère de choix concerne les performances : lorsque vous êtes connecté au pool de modems, vous arrivez en fait sur un réseau partagé par des centaines d'utilisateurs. Ce réseau est ensuite connecté à un autre, plus rapide, mais sur lequel il y a des milliers d'utilisateurs, et ainsi de suite. Or, pour des questions de rentabilité, les ISP ont intérêt à connecter le plus d'utilisateurs possible sur le réseau le moins rapide possible (car moins cher). Les ISP espèrent que leurs abonnés ne se connecteront pas tous en même temps : le ratio couramment admis dans la profession est d'un modem pour 30 ou 50 abonnés. Comme pour la voiture, évitez donc les heures de pointe !

Pour vous aider dans votre choix, des revues spécialisées présentent chaque mois un classement des ISP en fonction des performances et des pannes. Ce classement évolue sans cesse car les ISP procèdent continuellement à des réajustements de leur réseau.

De logiciels de communication

Un ordinateur n'est rien sans logiciel. Pour se connecter à l'Internet, il faut donc : un pilote de périphérique, une pile TCP/IP et un navigateur.

Pilote, TCP/IP, navigateur : voilà encore des termes abscons, mais qui feront bientôt partie de votre quotidien.

QUE VEUT DIRE " COUCHES RÉSEAUX " ?

Une communication s'établit à plusieurs niveaux : on se serre la main, on parle, on échange des regards. C'est la même chose en informatique : il y a plusieurs niveaux de communication : physique (signaux sur les câbles — le serrage de main), liaison (premier niveau indépendant du niveau physique), réseau (les voies de communication — les mots) et transport (le niveau de communication le plus élevé — les phrases). Au-dessus, il y a les applications (jeux, traitements de textes, etc.) et vous.

A chacun de ces niveaux est associée une couche logicielle numérotée de 1 à 4 ; on parle ainsi de couche transport, de niveau 2 pour la couche liaison.

À chaque niveau, l'information est structurée différemment : en **bits** pour le niveau physique, en **trames** pour le niveau liaison, en **paquets** pour les couches 3 et 4 (on parle aussi de datagramme pour la couche 3) et en **segments** pour les applications.

Les termes " niveau liaison ", " couche liaison ", " couche 2 ", " niveau 2 " et " niveau trame " sont équivalents.

Les pilotes de périphériques, couramment appelés *drivers* en anglais, ne sont quasiment jamais visibles pour les utilisateurs de base que nous sommes, mais ce sont eux que l'on installe et paramètre en premier.

Un *driver* permet à l'ordinateur de piloter un périphérique, c'est-à-dire un disque dur, un lecteur de CD-ROM ou... un modem, ce qui nous intéresse ici. Normalement, ils sont fournis avec ledit périphérique ou déjà intégrés dans Windows.

À un moment ou à un autre de la procédure d'installation, on vous demande toujours de quel modem vous disposez (à moins qu'il ne soit détecté automatiquement) et d'insérer la disquette ou le CD-ROM contenant le fichier X, généralement "X.inf".

Le deuxième logiciel nécessaire à une connexion Internet est TCP/IP : on parle de pile TCP/IP (*stack* TCP/IP, en anglais) ou encore de couches TCP/IP. Ce logiciel est inclus dans Windows et, s'il n'est pas installé, il vous sera demandé à un moment ou à un autre.

MAIS QU'EST-CE QUE TCP/IP ?

TCP/IP est le protocole (ou plutôt la famille de protocoles) utilisé sur l'Internet. **IP** (*Internet Protocol*) est le protocole de niveau 3 (couche réseau) ; **TCP** (*Transport Control Protocol*) est le protocole de niveau 4 (couche transport, comme son nom l'indique). Ce protocole permet aux ordinateurs d'échanger des informations découpées en **paquets**. Votre PC n'est pas relié directement au serveur situé à l'autre bout du monde (on pourrait le faire, mais ce serait très coûteux) ; vous devez donc passer par plusieurs intermédiaires (votre opérateur téléphonique préféré, votre ISP et d'autres encore). IP se charge d'acheminer ces paquets à travers ces différents intermédiaires (appelés nœuds du réseau) tout comme le fait la poste avec le courrier. C'est ce qu'on appelle du **routage**. Et, toujours comme avec La Poste, vous êtes identifié par une adresse : une **adresse IP** (ou encore une adresse réseau).

TCP est un protocole permettant d'engager une conversation de niveau supérieur, c'est-à-dire entre le client (vous, en France) et le serveur (situé à San Francisco, par exemple), et ce sans se soucier des intermédiaires. C'est en quelque sorte comme avec le téléphone : vous parlez directement à votre interlocuteur, sans vous soucier des machines et câbles intermédiaires qui véhiculent votre voix.

Enfin, un logiciel est nécessaire pour utiliser le réseau. Le plus répandu est le navigateur web (*browser Web*, en anglais). En anglais, web désigne la toile d'araignée. C'est ce qu'est en fait l'Internet : un gigantesque réseau d'ordinateurs reliés par des fils. Le navigateur permet d'y circuler sans danger, autrement dit de surfer sur le web.

Les navigateurs sont gratuits : on trouve en standard Internet Explorer, de Microsoft, et Communicator, de Netscape.

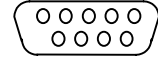
PC, Internet, ISP, driver, stack TCP/IP, browser web : voilà, vous connaissez désormais le jargon de base du réseau.

Comment faire fonctionner tout cela ?

Maintenant que vous disposez d'un port série (sur un PC, bien sûr), d'un modem, d'une ligne téléphonique, d'un abonnement à l'Internet et d'un navigateur, il faut assembler ce Mécano.

Identifier le port série sur le PC et y connecter le modem

Le port série est reconnaissable à la forme du connecteur DB9,

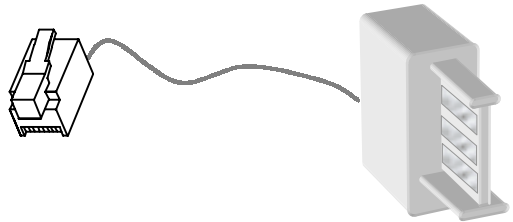


mâle côté PC et femelle côté modem. Sur le PC, il est également identifié par le symbole "O-O" affiché à côté. Le câble série doit donc être branché sur ces deux connecteurs.

Le modem doit ensuite être connecté à la prise téléphonique. Là encore, le câble est généralement fourni.

Il comporte d'un côté une prise RJ11 mâle qui ressemble à ceci :
et, à l'autre extrémité, une prise gigogne identique à celle de votre téléphone.

Il se peut, si vous êtes dans un hôtel ou au bureau, que la prise soit différente : on retrouve le format RJ11 aux États-Unis, ou un format analogue, appelé RJ45, tous deux étant du genre femelle.



Les prises ont la même forme, la RJ11 étant de taille plus petite. Cette dernière comporte quatre fils, généralement de couleurs noir, rouge, vert et jaune, tandis que la prise RJ45 comprend huit fils de couleurs bleu, orange, noir, rouge, vert, jaune, marron et gris. On peut apercevoir ces fils car les prises sont généralement transparentes.

Si vous avez à votre disposition une prise RJ45 femelle, vous pouvez néanmoins y connecter votre câble qui possède un connecteur RJ11 mâle (l'inverse n'est pas possible) : les formats sont, en effet, compatibles.

Côté modem, il faut brancher le câble sur le port RJ11, à côté duquel est affiché le mot "Line".

Si vous ne disposez que d'une ligne téléphonique et que vous vouliez conserver votre téléphone, il est possible de le chaîner au modem. Deux solutions : soit vous emboîtez les prises gigognes (celles du téléphone et du modem) au niveau de la prise téléphonique, soit vous connectez le téléphone sur la prise RJ11 intitulée "Phone" du modem.

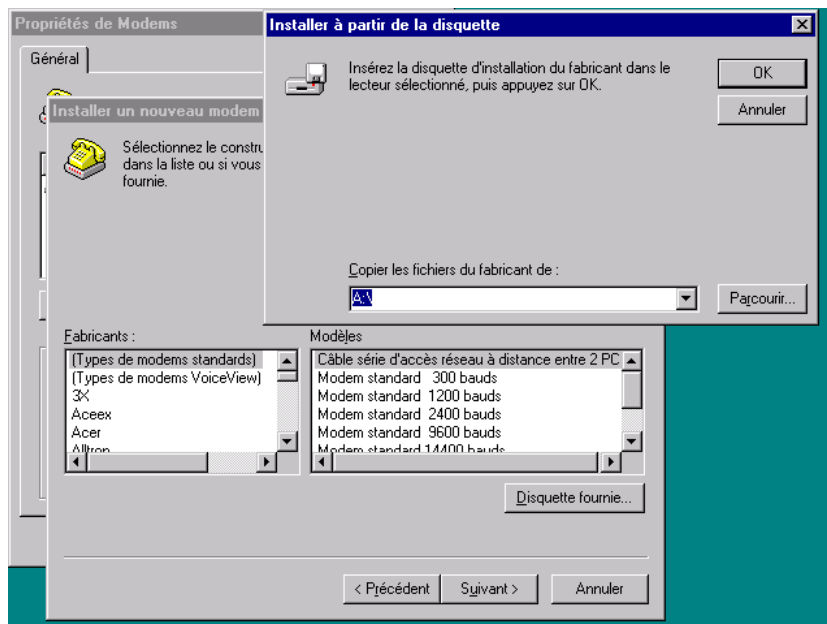
Installer et configurer le pilote du modem

Lorsque vous allumez votre PC, la fonction Plug and Play de Windows 9x détecte la présence d'un nouveau modem, et vous propose différents écrans de configuration. Cette étape correspond à l'installation d'un driver. Le driver est spécifique à chaque périphérique, ici notre modem. Il contient la liste des commandes permettant de piloter automatiquement le modem (paramétrage, numérotation, etc.).

Si le modem n'est pas détecté par la fonction *Plug and Play* de Windows, vous pouvez lancer manuellement la procédure en cliquant sur le menu "Démarrer→Paramètres→Panneau de configuration→Modems→Ajouter...".

Vous pouvez alors soit laisser le programme détecter automatiquement le type de modem à partir d'une liste standard incluse dans Windows, soit le sélectionner manuellement dans la même liste.

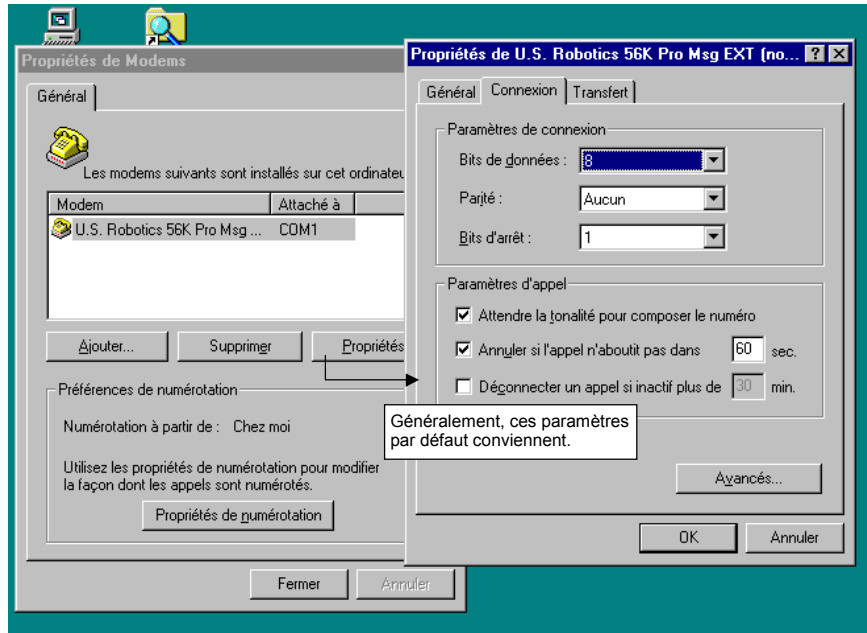
Vous obtenez alors l'écran suivant (quasiment identiques sous Windows NT et Windows 9x).



Si le modem est répandu et que son constructeur a passé des accords avec Microsoft, il figure dans la liste. Sinon, il faut cliquer sur "Disquette fournie...", puis insérer la disquette ou le CD-Rom accompagnant le modem dans son emballage.

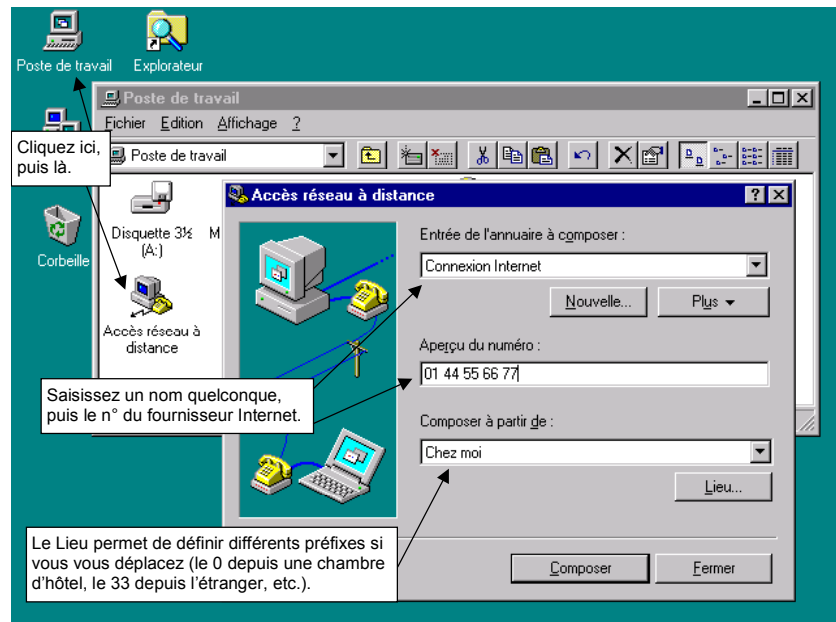
Indiquez ensuite le port série de l'ordinateur sur lequel est connecté le modem. En général, il s'agit du port COM1 ou COM2. Cliquez ensuite sur "Terminer", et réinitialisez l'ordinateur si cela est demandé.

Vous pouvez modifier à tout moment ces paramètres en cliquant sur le menu "Démarrer→Paramètres→Panneau de configuration→Modems→Propriétés".



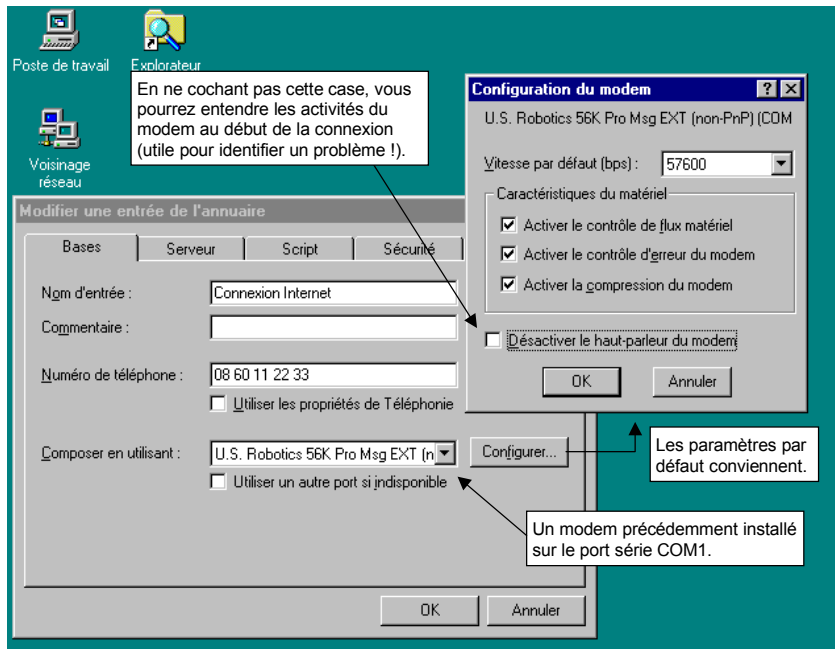
Configurer l'accès réseau à distance

Cette nouvelle étape consiste à définir les paramètres qui vont permettre à l'ordinateur d'utiliser le modem à l'aide du *driver*. Au niveau du bureau de Windows (écran standard qui apparaît après avoir démarré le PC), cliquez sur l'icône appelée "Poste de travail", ou "Mon ordinateur", puis cliquez sur "Accès réseau à distance".



L'exemple retenu ici est celui d'une connexion à un ISP quelconque (Wanadoo, Club-Internet, etc.).

Cliquez ensuite sur “Nouvelle...”, ou sur “Plus→Modifier l'entrée et les paramètres du modem...”, selon que vous procédez à l'opération pour la première fois ou que vous voulez modifier une précédente définition. Vous arrivez à l'écran suivant.



Saisissez le numéro de téléphone à appeler : il s'agit de celui du fournisseur d'accès Internet qui est indiqué sur la confirmation de l'abonnement auquel vous avez souscrit.

Le modem indiqué doit être celui préalablement installé. En cliquant sur “Configurer”, vous pouvez vous assurer que les paramètres par défaut conviennent. La valeur indiquée dans le champ “Vitesse par défaut” indique le débit, en bits par secondes, de la liaison entre votre modem et votre PC. Cela ne veut pas dire que le débit de ligne (entre votre modem et votre fournisseur Internet) sera identique. Généralement, on met une valeur égale ou légèrement supérieure à la vitesse maximale de ligne supportée par le modem (ici 56 Kbit/s).

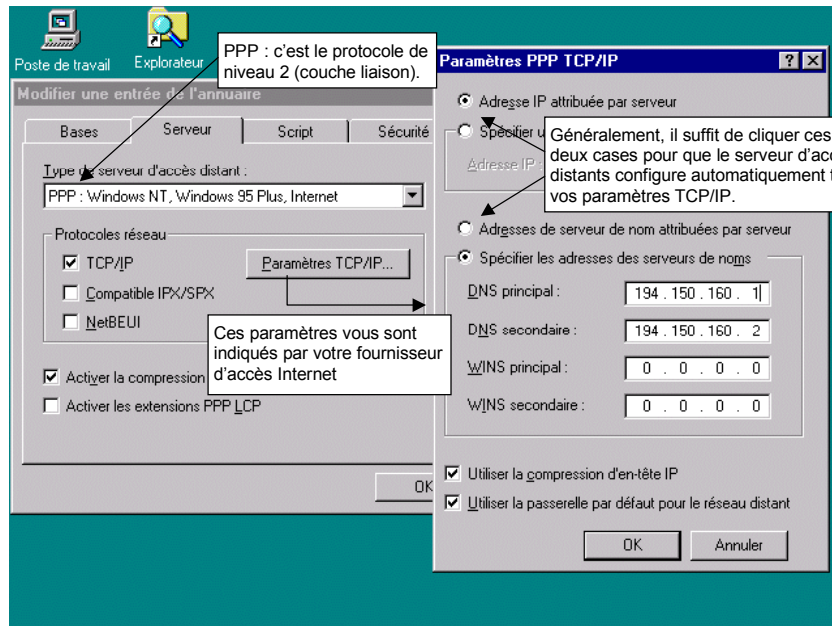
L'onglet suivant, intitulé “Serveur”, permet de configurer le protocole de niveau 2 (couche liaison) utilisé par le fournisseur d'accès Internet, généralement PPP (*Point-to-Point Protocol*).

Votre modem permet d'échanger des données avec celui du fournisseur d'accès : il gère la couche physique (niveau 1) en transmettant des signaux sur le réseau téléphonique commuté (RTC). Le protocole PPP, quant à lui, gère la couche liaison (niveau 2) : il permet à votre

ordinateur d'échanger des données avec le serveur de l'ISP (*Internet Service Provider*), votre fournisseur d'accès à l'Internet.

Le dialogue s'établit donc à plusieurs niveaux : physique (entre modems) puis logique (entre ordinateurs).

Le serveur est un équipement spécialisé, appelé **serveur d'accès distants**, qui gère un pool de modems. À chaque appel, ce serveur attribue un modem libre à la nouvelle demande de connexion et négocie des paramètres pour configurer les différents protocoles, dont TCP/IP.



Les paramètres essentiels de cet onglet concernent le protocole IP (*Internet Protocol*) qui gère la couche réseau (niveau 3). Alors que PPP ne gère qu'une liaison point à point (entre votre ordinateur et le serveur d'accès distant *via* la ligne téléphonique), IP élève le niveau de communication en assurant les échanges avec plusieurs nœuds du réseau.

Au sens réseau, un nœud est un équipement quelconque supportant un protocole de niveau 3 (ici IP). Et, tout comme un habitant d'une rue, un nœud est identifié par une adresse réseau, ici une adresse IP.

Dans notre exemple, l'adresse est attribuée par le serveur lorsque vous vous connectez à l'Internet. Cela permet au fournisseur d'accès de s'assurer que votre ordinateur (désormais un nœud du réseau Internet !) aura bien une adresse IP unique sur le réseau.

Si vous spécifiez une adresse IP attribuée par vous-même, il y a de fortes chances pour que vous ne soyez pas reconnu sur l'Internet. Avec une fausse adresse ou une adresse identique à

QU'EST-CE QU'UNE ADRESSE IP ?

Une adresse permet d'identifier de manière unique un **nœud** connecté à un réseau. Un nœud est un terme générique désignant un ordinateur ou tout autre équipement connecté à un réseau.

Lorsque vous recevez un paquet postal, celui-ci comporte l'adresse du destinataire (vous) et celle de l'émetteur. Il en est de même pour un paquet IP. Les adresses permettent d'acheminer les paquets IP à travers l'Internet jusqu'au destinataire.

L'Internet impose un adressage qui doit être respecté par tous. On parle d'**adresses publiques**, car l'Internet est un réseau public, tout comme le sont les rues d'une ville. Les adresses publiques sont attribuées par des organismes officiels qui gèrent l'**espace d'adressage** de l'Internet, tout comme la mairie affecte les adresses aux habitants.

Dans une ville, il y a des rues et des numéros d'immeubles dans chaque rue. De même, sur l'Internet, il y a des réseaux et des numéros de nœuds dans chaque réseau.

Le format standard d'une adresse Internet est constitué de quatre chiffres décimaux séparés par des points, compris entre 0 et 255, par exemple **194.160.150.2**. Une partie des chiffres désigne le numéro de réseau (unique au sein de l'Internet), l'autre le numéro du nœud au sein de ce réseau (unique au sein du réseau considéré).

celle d'une autre personne, votre courrier postal n'arrivera jamais, ou alors aléatoirement. Sur l'Internet, c'est pareil.

Tester votre connexion

Vous venez de relier physiquement le modem à l'ordinateur (*via* un câble) et vous avez relié logiquement le modem au logiciel TCP/IP (*via* un driver) en configurant l'accès réseau à distance. Pour tester la connexion, il suffit de cliquer, au niveau du bureau, sur l'icône appelée "Poste de travail", ou "Mon ordinateur", puis sur l'icône "Accès réseau à distance". Cliquez enfin sur "Composer".

Que se passe-t-il alors ?

L'ordinateur pilote automatiquement le modem : cela est confirmé lorsque vous entendez la numérotation des dix chiffres. Vous entendez ensuite une succession de bruits différents : tout d'abord, la sonnerie du téléphone, jusqu'à ce qu'un modem décroche du côté de l'opérateur (le fournisseur d'accès à l'Internet, l'ISP), puis un sifflement aigu qui indique la prise de ligne (c'est ce qu'on appelle la porteuse).

Vous entendez alors un bruit ressemblant à un "dong" ou un "dong-dong" : cela signifie que les modems sont en train de négocier une vitesse à partir de la plus haute possible (dans notre exemple, 56 Kbit/s). Si le modem de l'opérateur ne supporte que 28,8 Kbit/s, votre modem rétrogradera à cette vitesse. Si la qualité de ligne est mauvaise (détectée par des erreurs de transmission), les deux modems rétrograderont jusqu'à trouver une bonne qualité de ligne. C'est ce qui s'est passé dans l'exemple ci-après.



Puis, vous entendez un chuintement qui indique que des données sont en train d'être échangées : vos nom de compte et mot de passe sont envoyés au serveur de l'opérateur qui les vérifie par rapport sa base de données, puis renvoie une autorisation. Les paramètres des différentes couches réseau, PPP et TCP/IP, sont également négociés. Le serveur vous attribue notamment une adresse IP.

Enfin, le message "Connecté à : 28,8 Kbit/s" apparaît en bas de l'écran : vous êtes connecté à l'Internet !

Résoudre les problèmes de connexion

Si, comme moi vous n'êtes pas satisfait de la vitesse de 28,8 Kbit/s, vous pouvez essayer une nouvelle connexion jusqu'à obtenir la vitesse pour laquelle vous avez payé votre abonnement. Si, au bout de deux ou trois tentatives, vous n'obtenez toujours pas la bonne vitesse, c'est sans doute que les "bons" modems sont déjà pris par d'autres utilisateurs : bienvenue dans la jungle Internet !

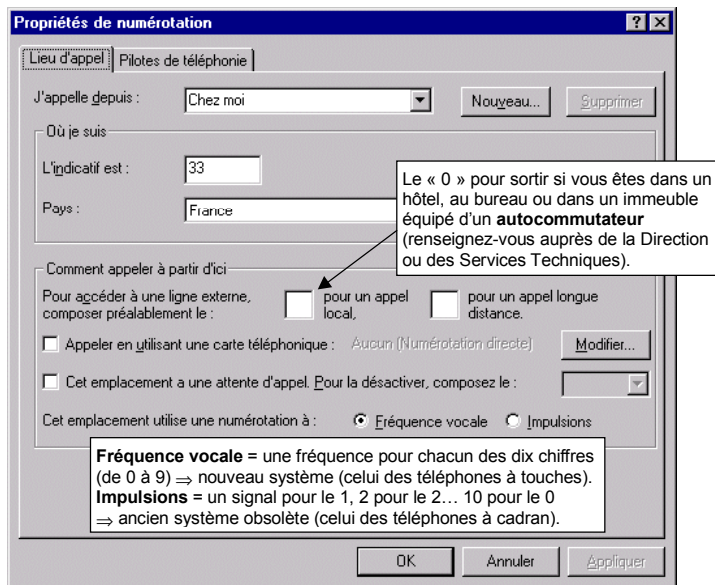
Essayez de nouveau en période creuse (le matin, par exemple). Si vous n'obtenez toujours pas la vitesse maximale autorisée par votre modem, c'est qu'il y a un problème avec votre ligne téléphonique ou que vous appelez un numéro sur lequel il n'y a que des "vieux" modems. Vérifiez auprès de votre fournisseur d'accès que le numéro de téléphone correspond bien à des modems 56 Kbit/s ou plus.

Plus grave : rien ne fonctionne !

➔ Le modem ne numérote pas. Vérifiez que tout est en ordre côté matériel : modem allumé, câbles connectés (modem→PC et modem→prise téléphonique). Vérifiez ensuite

que vous êtes connecté au même port série que celui sélectionné lors de l'installation du modem (COM1 dans notre exemple).

- ➔ Le modem numérote, mais personne ne répond, ou vous tombez sur un faux numéro. Si vous êtes dans un hôtel ou sur votre lieu de travail, vous devez ajouter un préfixe pour l'extérieur, généralement le "0". Si vous êtes chez vous, il n'en faut pas. Cliquez sur le menu "Démarrer→Panneau de configuration→Téléphonie", et vérifiez les paramètres pour votre lieu d'appel (dans notre exemple, "Chez moi").



Si, malgré tout, vous tombez systématiquement sur un faux numéro, vérifiez auprès de votre fournisseur d'accès que ses modems sont compatibles avec le vôtre.

Si vous êtes derrière un autocom privé (raccourci pour dire que votre ligne téléphonique est raccordée à un autocommutateur de votre entreprise ou de l'hôtel — sur une carte analogique, bien sûr), essayez avec une ligne directe, c'est-à-dire qui ne passe pas par l'autocommutateur.

Si vous êtes relié à un autocommutateur et que vous appelez un modem également raccordé à un autre autocommutateur, il se peut qu'il y ait une incompatibilité entre eux : votre modem numérote, la sonnerie retentit à l'autre bout, puis plus rien. Vous n'entendez même pas le modem distant décrocher. Les autocommutateurs sont généralement installés sur des réseaux numériques et peuvent échanger des informations de manière incompatible. Il n'y a alors rien à faire : l'un des deux correspondants doit utiliser une ligne analogique directe.

- ➔ Le modem de l'opérateur ne décroche pas ou sonne occupé. Dans ce cas, c'est de lui que vient le problème. La seule chose à faire est alors d'appeler son service de "support client".

Les applications du réseau : l'exemple du web

Ça y est : vous êtes sur l'Internet ! Mais que faire ? Par où commencer ? L'Internet offre, en effet, une multitude de services, tous reposant sur le même principe : la connexion à des serveurs web qui permettent d'afficher des pages à l'écran et de télécharger des fichiers. Ces pages sont riches en couleurs, en images, en sons, en vidéos et, accessoirement, en écrans publicitaires. Les fichiers peuvent, quant à eux, être de nature différente : logiciels, pilotes, documentation, musique, images, vidéo, etc.

Toutes les opérations complexes permettant d'y accéder sont masquées par un logiciel qui a révolutionné le monde informatique : le navigateur web.

L'Internet permet également d'échanger des informations à travers des forums de discussion (les *news*) et la messagerie — qui permet d'envoyer et de recevoir les fameux e-mails (*electronic mail*).

Ce chapitre vous apprend donc comment utiliser le réseau Internet, c'est-à-dire :

- à configurer et à utiliser un navigateur ;
- les concepts de base du web ;
- des notions sur HTTP, HTML et les URL ;
- les principes du transfert de fichiers FTP ;
- à configurer et à utiliser la messagerie ;
- et, en définitive, ce qu'est l'Internet.

Configurer et utiliser un navigateur

La première application que l'on installe généralement est un navigateur (appelé également *browser web*) : il s'agit d'Opera, de Microsoft Internet Explorer ou de Netscape Navigator (ou encore de Netscape Communicator qui inclut la messagerie). Nous retiendrons ce dernier pour nos exemples.

Qu'est-ce que le web ?

Le web (la toile d'araignée en anglais) est constitué de serveurs, appelés serveurs web, auxquels vous pouvez vous connecter à l'aide d'un navigateur. Un serveur web est semblable à un serveur Vidéotex, et le navigateur web peut être assimilé à un Minitel. La différence majeure est que le web est multimédia : vous pouvez visualiser des images (fixes ou animées) et des films ou écouter de la musique, alors que le Minitel affiche péniblement des caractères semi-graphiques vaguement animés, à la vitesse record de 1,2 Kbit/s. Autre différence : le Minitel repose sur des protocoles propriétaires tandis que ceux utilisés sur l'Internet sont des standards.

Le nommage des pages web

Nous l'avons vu au chapitre précédent, « Configurer l'accès réseau à distance », les nœuds du réseau Internet (votre PC et les serveurs) sont identifiés par des adresses IP. Afin d'éviter aux utilisateurs d'avoir à retenir des dizaines d'adresses numériques du type 196.129.214.159, l'Internet fournit un service de noms appelé DNS (*Domain Name System*) : les utilisateurs saisissent le nom du serveur sur lequel ils veulent se connecter, puis le système convertit ce nom en une adresse IP.

QU'EST-CE QUE LE NOMMAGE INTERNET ?

Au chapitre précédent, nous avons vu que chaque ordinateur connecté à l'Internet dispose d'une adresse réseau unique (une adresse IP). Le problème est que cette adresse n'est connue de personne : *a priori*, vous ne connaissez pas l'adresse du serveur web de la société X située en Californie. De plus, ce serveur peut déménager et donc changer d'adresse IP.

Pour faciliter la vie des internautes, un **espace de nommage** a donc été superposé à l'espace d'adressage IP. Le principe est le même que pour les adresses postales : dans chaque ville il y a des noms de rues qui sont d'ailleurs souvent identiques. Quelle ville ne possède pas son Avenue du Général De Gaulle ? Sur l'Internet, il y a des **domaines** (COM pour commercial, FR pour France, EDU pour éducation, etc.), dans chaque domaine, des noms de sociétés (Laposte, edf, cisco, etc.), puis dans chaque société, des serveurs web appelés par convention **www** (*World Wide Web*).

Il existe donc au sein de l'Internet des **serveurs de noms** qui gèrent des bases de données contenant les correspondances entre adresses IP et noms. L'ensemble de ces serveurs constitue le **DNS** (*Domain Name System*). Lorsque vous saisissez `www.laposte.fr`, votre navigateur web interroge le serveur de noms le plus proche (généralement celui de votre ISP) pour lui demander l'adresse IP du serveur français de La Poste, appelé `www`. Ce mécanisme s'appelle la **résolution de noms**.

Sur l'Internet, vous accédez au service de noms *via* votre ISP. Le DNS est expliqué en détail au chapitre 17.

Naviguer

Nous avons choisi ici Netscape comme exemple, mais les principes et l'affichage sont les mêmes avec les autres navigateurs, tels que Internet Explorer et Opera.

La navigation consiste à indiquer le nom d'un serveur au format URL, par exemple `www.iana.org` ou `myweb.vector.ch`. Le nom du serveur est, par convention, `www` (sigle de *World Wide Web*), mais il peut être choisi librement, par exemple `myweb`. Le nom de domaine DNS, dans notre exemple “ `vector.ch` ”, doit être déposé auprès des organismes officiels de l'Internet (voir la fin de ce chapitre à ce sujet).

LE POINT SUR LES URL (RFC 1630,1738)

Les **URL** (*Uniform Resource Locator*) décrivent la manière d'accéder à des informations sur l'Internet, c'est-à-dire le protocole ainsi que la syntaxe utilisés pour accéder à des fichiers.

La syntaxe générale d'une URL est **<protocole>://<syntaxe spécifique au protocole>**. Elle est reconnue par tous les navigateurs web et par de plus en plus de logiciels tels que les traitements de texte.

Parmi les valeurs du champ **<protocole>** reconnues par les navigateurs web, on trouve :

- **ftp** (*File Transfer Protocol*) pour les transferts de fichiers ;
- **http** (*HyperText Transfer Protocol*) pour le web ;
- **mailto** pour les adresses de messagerie électronique (*e-mail*) ;
- **nntp** (*Network News Transfer Protocol*) pour les forums de discussion ;
- **telnet** pour les connexions en émulation de terminal ;
- **file** pour l'accès au système de fichiers du disque dur local ;

ainsi que : **gopher**, **news**, **wais** (*Wide Area Information Server*) et **prospero** (service d'annuaire).

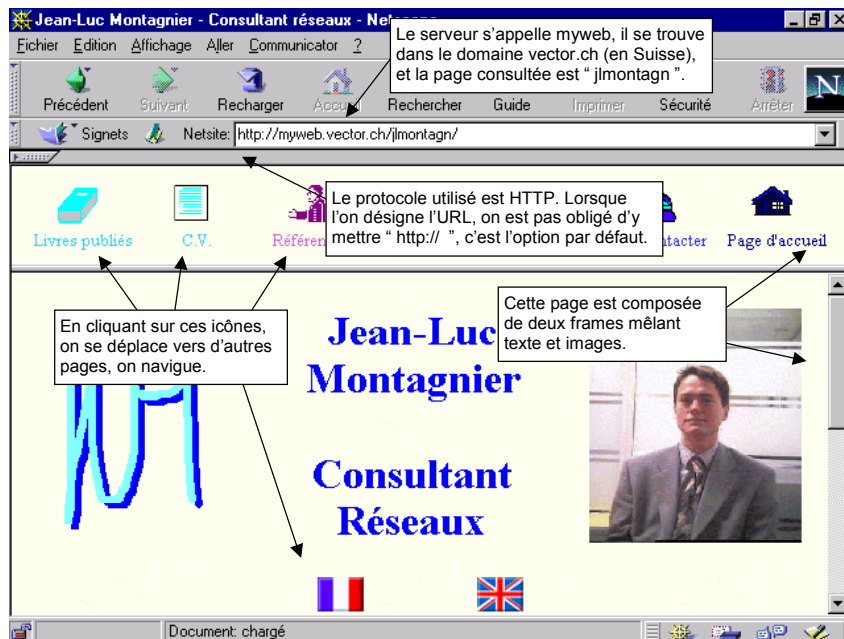
Tous ces protocoles sont dits de niveau applicatif ; ils sont situés au-dessus de la couche TCP.

Voici les syntaxes spécifiques à quelques protocoles :

- `ftp://<compte utilisateur>:<mot de passe>@<machine>:<port TCP>/<chemin d'accès au fichier>; type=<code>` (<code> désigne une commande FTP : `a=ascii`, `i=image`, `d=dir`)
- `http://<machine>:<port TCP>/<chemin>?<recherche>`
- `mailto:<adresse e-mail>`
- `nntp://<machine>:<port TCP>/<nom du groupe de news>/<n° de l'article>`
- `telnet://<compte utilisateur>:<mot de passe>@<machine>:<port TCP>/`
- `file://<chemin d'accès au système de fichiers>` (par exemple, `file://C:/temp/` affiche le contenu du répertoire TEMP sur le disque dur C:).

De nombreux champs sont optionnels ; par exemple, si le port n'est pas spécifié, la valeur par défaut est celle du port standard du protocole (23 pour telnet, 80 pour HTTP, etc.). Dans le cas de FTP, si les champs **<user>** et **<password>** sont omis, le compte par défaut est : `anonymous`.

Remarque Le champ **<machine>** désigne l'adresse IP ou le nom DNS du serveur cible.



L'affichage d'une telle page peut sembler complexe, mais elle consiste, pour le navigateur, à interpréter un fichier texte au format ASCII contenant des commandes HTML (*HyperText Markup Language*).

Par exemple, le fichier suivant contient les instructions permettant au navigateur d'afficher l'écran présenté ci-dessus.

```
<HTML>
<HEAD>
  <META HTTP-EQUIV="Content-Type"
  CONTENT="text/html; charset=iso-
  8859-1">
```

```
  <META NAME="description" CONTENT="Jean-
  Luc Montagnier - Consultant réseaux">
```

QU'EST CE QUE HTTP ET HTML ?

Le protocole **HTTP** (*HyperText Transfer Protocol*) permet à un navigateur d'interroger un serveur web dans le but de télécharger des fichiers vers un PC. Ces fichiers doivent être au format HTML (*HyperText Markup Language*), qui est un langage de programmation ou plutôt de description d'une **page web**. Un fichier HTML peut contenir des instructions telles que le téléchargement d'autres **objets** qui composent la page web (images fixes ou animées, son, vidéo, etc.) ou l'exécution de **scripts** et d'**applets Java** (ce sont de petits morceaux de programmes téléchargés à la demande).

En-tête standard permettant à la page d'être référencée dans les moteurs de recherche.

```

<META NAME="keywords" CONTENT="Montagnier, Consultant, Indépendant,
consulting, informatique, réseau, réseaux, service, services">
<META NAME="GENERATOR" CONTENT="Mozilla/4.02 [en] (WinNT; I) [Net-
scape]">
<TITLE>Jean-Luc Montagnier - Consultant r&eacute;seaux</TITLE>

</HEAD>
<BODY BACKGROUND="x5.jpg">
  &nbsp;
  <CENTER><TABLE BORDER=0 WIDTH="97%" >
  <TR>
  <TD ALIGN=Center VALIGN=Center WIDTH="165">
  <CENTER><IMG SRC="acc-2.gif" VSPACE=10 HEIGHT=160 WIDTH=106></CENTER>
  </TD>
  ...
  <CENTER><A HREF="Us-home.htm" TARGET="_parent"><IMG SRC="Drap-uk.gif"
  BORDER=0 HEIGHT=26 WIDTH=35></A></CENTER>

  <CENTER>English version</CENTER>
  ...

```

Indique de charger ces fichiers et de les afficher sous forme d'images (formats JPEG et GIF dans notre exemple).

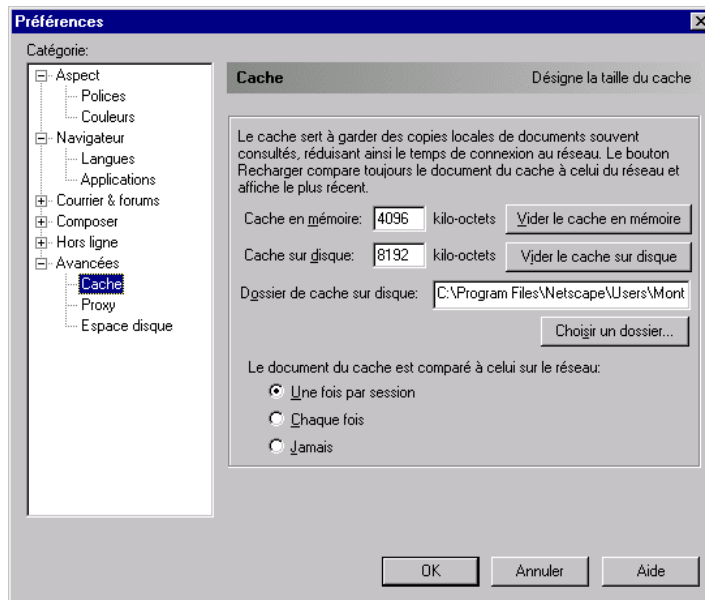
Affiche un lien URL sous forme d'icône (fichier "Drap-uk.gif"). En cliquant dessus, vous vous déplacez vers la page "Us-home.htm".

Il est possible de composer une page web de ce type manuellement à l'aide d'un éditeur de texte, mais il est plus simple d'utiliser un composeur de pages (tel que celui inclus dans le navigateur de Netscape) qui génère automatiquement le fichier texte. La composition d'une page consiste alors à faire glisser à l'aide de la souris les éléments que l'on veut afficher (frames, images, texte, liens URL, etc.).

En résumé, une page web est un simple fichier que le navigateur télécharge depuis un serveur situé sur l'Internet. Elle peut comporter des ordres permettant de télécharger d'autres fichiers tels que des images. Plus la page est complexe, plus sa taille sera importante et plus son chargement sera lent. En outre, les pages peuvent contenir des animations, des sons, etc.

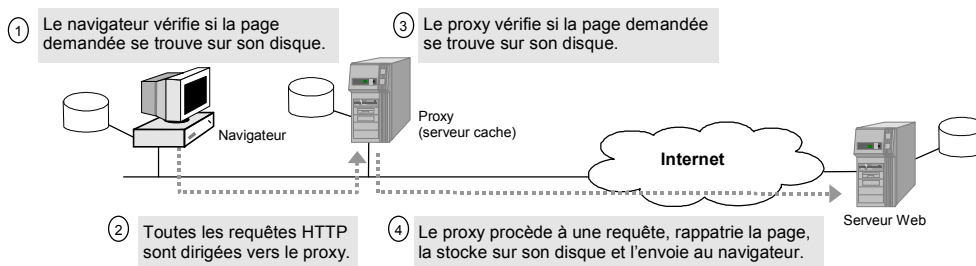
Améliorer les performances

Lorsque vous naviguez sur le web, vous passez sans doute plusieurs fois sur la même page, ou encore les mêmes logos et images peuvent s'afficher. D'un jour à l'autre, vous pouvez retourner sur le même site. Afin d'éviter de recharger ces mêmes informations, et donc pour gagner du temps, le navigateur peut stocker localement ces pages sous forme de fichiers et gérer ce qu'on appelle un **cache**. Avec Netscape, elles sont stockées dans le répertoire \Program Files\Netscape\Users\nom_utilisateur\Cache. Lorsque le même élément devra être affiché, le navigateur chargera le fichier précédemment enregistré au lieu de le télécharger à travers le réseau. Cette technique peut vous faire gagner du temps, surtout si vous êtes connecté *via* un modem.



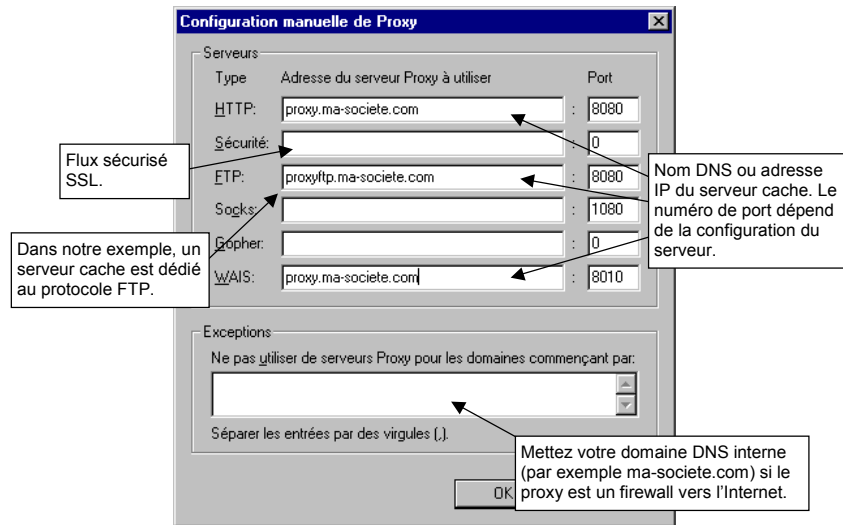
En plus du cache local à votre PC, il peut être intéressant d'utiliser un serveur cache, également appelé proxy cache, qui gère un cache pour tous les utilisateurs d'une entreprise (voir figure 3-3). Typiquement, un proxy est positionné entre la connexion Internet et le réseau interne à votre entreprise (votre intranet).

Figure 3-1.
Utilisation d'un serveur cache.



Attention à ne pas faire la confusion : le terme proxy est également utilisé pour désigner un certain type de firewall (ou pare-feu) qui n'a rien à voir avec un serveur cache. Cependant, si un tel firewall est installé entre votre réseau et l'Internet, vous devrez le désigner de la même manière qu'un proxy : dans le navigateur, le même paramétrage a en fait été étendu aux firewall.

À partir de l'écran précédent, cliquez sur " Proxy ", sélectionnez " Configuration manuelle du Proxy ", puis cliquez sur " Afficher... ". Vous obtenez alors l'écran suivant.



En activant cette fonction, toutes les requêtes web, WAIS (serveur d'archive) et FTP seront dirigées vers les serveurs proxy désignés.

Les cookies

Lorsque vous naviguez sur l'Internet, vos faits et gestes peuvent être surveillés, et des informations peuvent même être modifiées sur votre PC. Le but peut être commercial (il s'agit d'identifier les meilleurs clients potentiels) ou sécuritaire (recherche de malfaiteurs).

Un cookie (littéralement un petit malin, ou un biscuit, comme vous voulez !) est une information laissée par un serveur web sur votre poste de travail. Le navigateur de Netscape les stocke dans le fichier cookies.txt situé dans le répertoire \Program Files\Netscape\Users\nom_utilisateur.

```
# Netscape HTTP Cookie File#
http://www.netscape.com/newsref/std/cookie_spec.html# This is a generated file!
Do not edit.
```

www.selection.francetelecom.fr	FALSE	/SFT	FALSE	944035235
SHOPPERMANAGER%2FSFT				
JGK42FJ34PSH2PFN00CG1EAN6CSUE48R.microsoft.com	TRUE	/		
FALSE 937422045 MC1				GUID=B39F4640694311D19D3F0000F84121EB

Deux cookies ont été installés, un par le serveur Web de France Télécom et un autre par celui de Microsoft...

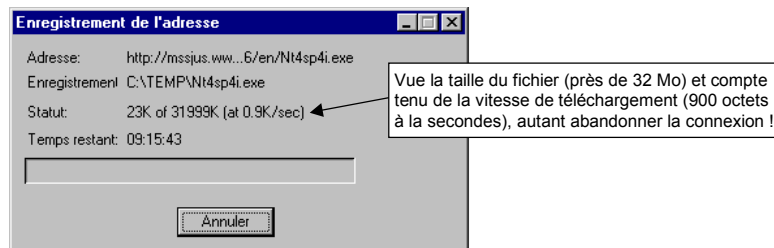
Le cookie aide le serveur web à savoir combien d'utilisateurs sont connectés (un utilisateur peut, en effet, effectuer de nombreuses requêtes) ou permet de conserver des informations sur votre profil : est-ce vous êtes déjà passé par cette page, quelles sont les services que vous

avez déjà consulté, est-ce que vous avez déjà effectué des achats, ou encore, quelle est la configuration de votre PC. C'est ainsi que vous pouvez voir s'afficher une publicité qui, oh ! hasard, correspond exactement à ce que vous recherchez...

Transférer des fichiers

Un autre programme fréquemment utilisé est FTP (*File Transfer Protocol*). En fait, vous l'utilisez sans vous en rendre compte chaque fois que vous rapatriez un fichier sur votre PC. C'est ce qui se produit lorsque le navigateur vous demande d'enregistrer un fichier.

L'écran suivant indique ensuite la progression du transfert depuis le serveur vers votre PC.



Inversement, vous pouvez envoyer des fichiers vers un serveur, ce qui peut être utile pour mettre à jour votre site web. La syntaxe de l'URL est la suivante : nom_de_l'utilisateur:mot_de_passe@myweb.vector.ch. Si le compte utilisateur et le mot de passe sont omis, vous vous connectez en mode anonyme (le compte "anonymous" est un compte invité avec des droits restreints).

Une fois connecté, vous pouvez, à l'aide de la souris, déplacer les fichiers depuis l'Explorateur vers la fenêtre du navigateur. Celui-ci vous demandera alors une confirmation et transférera les fichiers. L'opération inverse est possible en appuyant sur la touche « Maj. » (ou « ⌘ ») lors de la sélection du fichier à l'aide de la souris.

Les possibilités du navigateur sont cependant limitées à ces fonctions : vous ne pouvez pas, par exemple, supprimer des fichiers. Le mieux est de vous procurer un client FTP du commerce ou du domaine public, ou encore d'utiliser celui de Windows NT en mode texte.

QU'EST-CE QUE FTP ?

Le protocole FTP (*File Transfer Protocol*) est associé à l'application du même nom qui permet de transférer des fichiers au-dessus de TCP/IP. Ce protocole est très simple mais n'offre aucune **reprise sur erreur**. Si un problème réseau survient lors du transfert du dernier octet du fichier, il faudra recommencer la procédure depuis le début.

```

Invite de commandes - ftp myweb.vector.ch

C:\>ftp myweb.vector.ch
Connected to myweb.vector.ch.
220- UECTOR communication FTP server WAR-FTPD 1.65 Ready
220 Please enter your user name.
User <myweb.vector.ch:(none)>: xxxxxx
331 User name okay, Need password.
Password:
230-
230- Welcome on the UECTOR communication FTP server.
230-
230- Currently 4 user(s) is/are connected.
230-
230- If you have any problem using this service, please contact ftp@vector.ch.
230 User logged in, proceed.
ftp> bin
200 Type is Image <Binary>
ftp> put toto.exe
200 Port command okay.
150 Ready to receive "/E/ftp/MyWEB/jlmontagn/toto.exe". Mode STREAM Type BINARY

226- Transfer finished successfully. Closing data connection.
226 UL: 101 DL: 3 Ratio: 0:23924 Credit: unlimited *disabled
26800 bytes sent in 11,28 seconds <2,38 Kbytes/sec>
ftp>

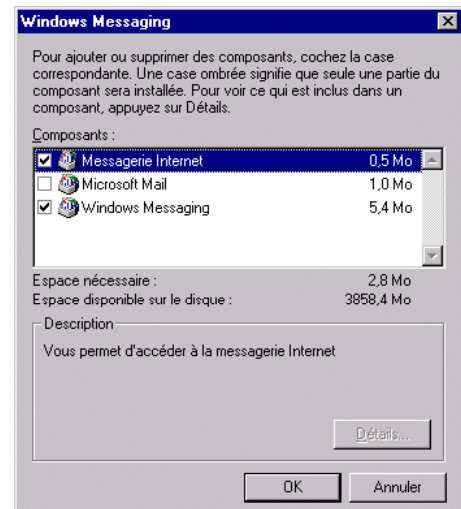
```

La commande “get” permet de réaliser le transfert inverse, depuis le disque du serveur vers le disque local. La commande “ascii” permet de revenir au mode de transfert par défaut qui consiste à convertir les codes ASCII des fichiers texte en fonction du type d’ordinateur employé.

Installer et configurer la messagerie

La troisième grande utilisation de l’Internet est la messagerie électronique (couramment appelée *e-mail*).

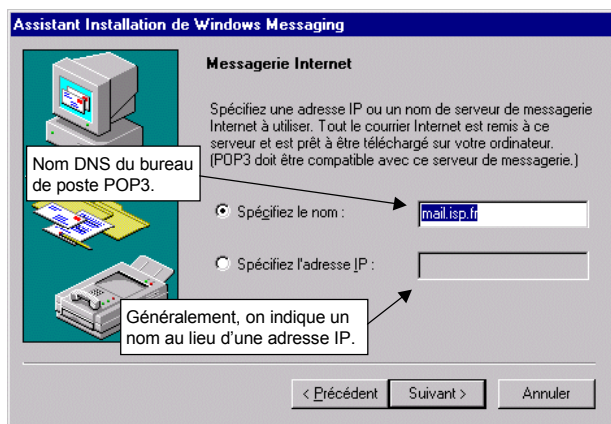
Internet Explorer ne dispose pas de messagerie intégrée. Il faut pour cela installer et configurer un autre programme, inclus dans Windows, qui s’appelle, selon les cas, Windows Messaging, Exchange ou Outlook Express. Pour cela, cliquez sur le menu “Démarrer→Paramètres→Panneau de configuration→Ajout/Suppression de programmes→Installation de Windows→Windows Messaging (ou Exchange)”, et vérifiez que les composants suivants sont installés.



Dans certaines versions de Windows 9.x, un menu similaire est accessible en cliquant sur l'icône " Boîte de réception " située sur le bureau. La première fois que vous cliquez dessus, le programme d'installation vous guide dans les étapes de configuration.

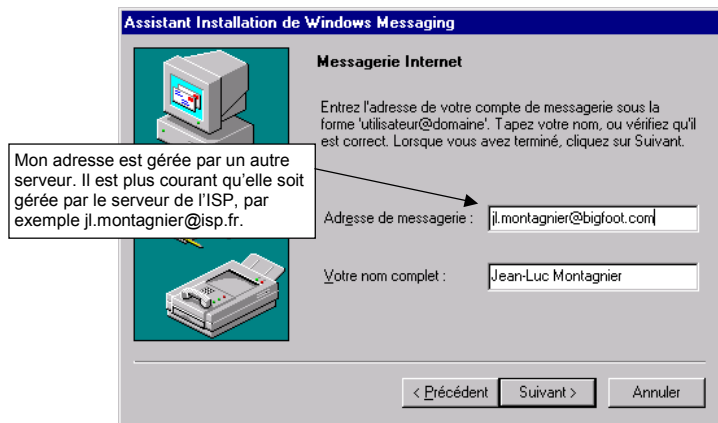
Vérifiez que les cases " Utiliser les services d'information suivants " et " Messagerie Internet " sont cochées, puis cliquez sur " Suivant ". Sélectionnez alors " Modem ", puis cliquez de nouveau sur " Suivant ".

Vous voyez alors apparaître la configuration que vous avez créée au cours du chapitre précédent, dans notre exemple, " Connexion Internet ". Vous pouvez en créer une nouvelle en cliquant sur " Nouveau " : vous retournez alors aux menus décrits à la section " Configurer l'accès réseau à distance " du chapitre précédent. Si vous cliquez sur " Suivant ", le programme vous demande de saisir le nom d'un serveur de messagerie.



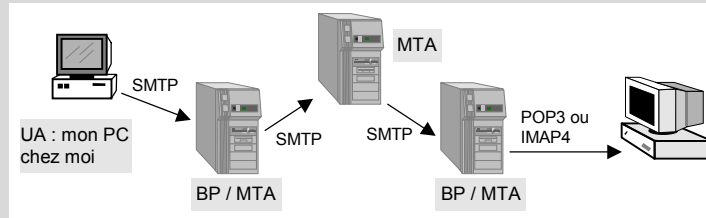
Il s'agit du serveur Internet qui va vous permettre d'envoyer et de recevoir du courrier électronique. Il faut ici saisir le nom communiqué par votre ISP, généralement mail.isp.fr ou pop.isp.fr.

Cliquez ensuite sur " Suivant ", assurez-vous que l'option " Autonome " ou l'option " Déconnecté " est activée, puis cliquez de nouveau sur " Suivant ". Un autre écran important apparaît.



COMMENT FONCTIONNE LA MESSAGERIE SUR INTERNET ?

Un système de messagerie repose sur des serveurs, appelés **MTA** (*Message Transfer Agent*) et des clients, appelés **UA** (*User Agent*). Les MTA ont pour tâche d'acheminer le courrier au sein de l'Internet et de les remettre aux UA destinataires. Un UA est rattaché à un MTA qui héberge des boîtes aux lettres (**BAL**, en abrégé) ; ce MTA fait alors office de **bureau de poste** (BP). Le client dépose des messages sur son bureau de poste.



Les messages sont transférés depuis le bureau de poste vers le PC à l'aide des protocoles **POP 3** (*Post Office Protocol*) ou **IMAP 4** (*Internet Message Access Protocol*). Les messages sont transférés depuis le PC vers le bureau de poste à l'aide du protocole **SMTP** (*Simple Mail Transfer Protocol*). Les MTA acheminent également les messages à l'aide de SMTP.

Il faut ici saisir votre adresse e-mail, c'est-à-dire votre adresse de messagerie qui sera connue de toutes les personnes qui voudront vous envoyer un message. Après avoir cliqué sur “ Suivant ”, le programme vous demande de saisir un autre nom.

Il s'agit cette fois d'un compte utilisateur correspondant à votre boîte aux lettres sur le bureau de poste (dans notre exemple, *mail.isp.fr*). Ce compte personnel vous permettra de vous connecter au serveur, pour y déposer et y retirer des messages. Le nom et le mot de passe sont directement utilisés par les protocoles POP3 et IMAP4.

QU'EST-CE QU'UNE ADRESSE DE MESSAGERIE ?

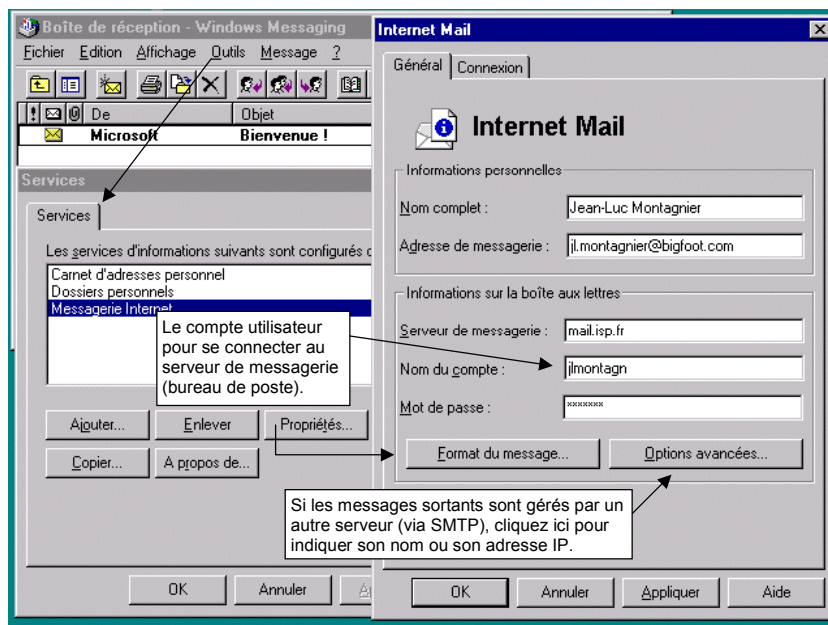
Une adresse de messagerie (**e-mail** pour *electronic mail*) permet de vous identifier de manière unique sur l'Internet. Tout comme votre ordinateur est identifié par une adresse IP et un nom au niveau de la couche réseau (couche IP), un utilisateur est identifié par un nom au niveau de l'application (ici la messagerie). Cet adressage se situe donc à un niveau supérieur et ne fait plus référence aux couches réseaux (1. physique modem, 2. liaison PPP, 3. réseau IP et 4. transport TCP).

Une adresse e-mail se présente sous la forme **nom@nom_de_domaine**, par exemple *jl.montagnier@bigfoot.com*. L'arobase, « @ », se prononce « at ». Le nom de domaine fait référence au **DNS**, l'espace de nommage Internet. Attention, ce nom est différent de celui du compte utilisé pour se connecter au bureau de poste. Ce dernier n'a qu'une portée locale : il n'est connu que du serveur (et de l'utilisateur). Vous devez par contre diffuser votre adresse e-mail le plus largement possible pour que l'on vous écrive.



Cliquez ensuite sur “ Suivant ”, puis validez le nom proposé pour votre carnet d’adresses personnel. Il s’agit du fichier qui contiendra la liste des adresses e-mail de vos correspondants. Cliquez de nouveau sur “ suivant ”, puis validez le nom proposé : il s’agit, cette fois, du fichier qui contiendra tous les messages stockés en local (ceux qui sont en attente d’envoi et ceux que vous avez rapatriés depuis le serveur de messagerie). Cliquez enfin sur “ Terminer ”.

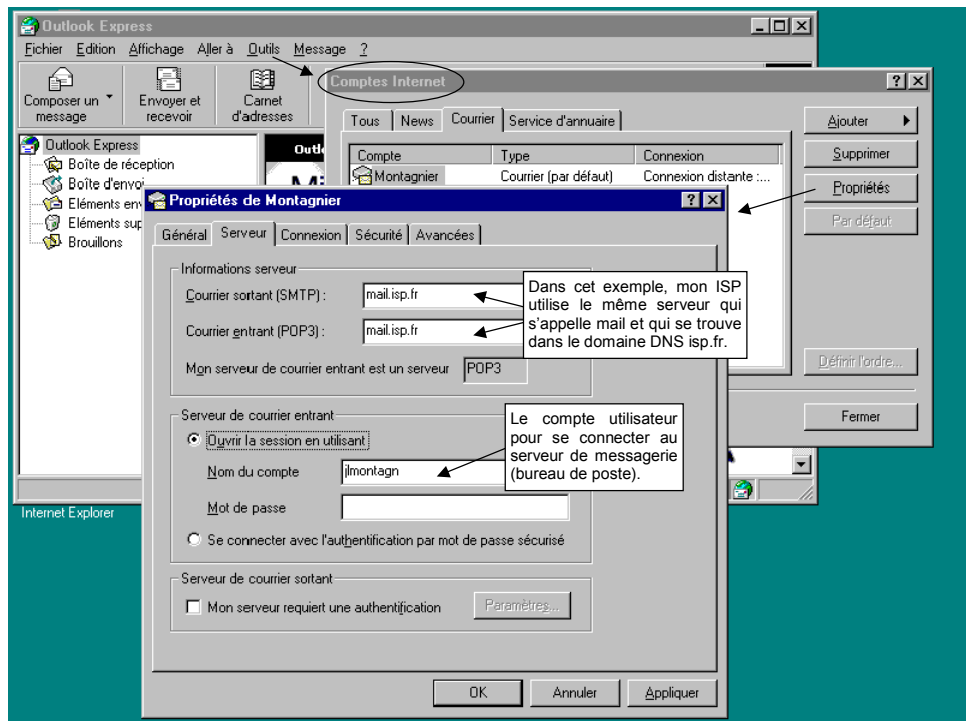
Il est possible de modifier tous ces paramètres à partir du menu “ Outils→Services→Propriétés ”.



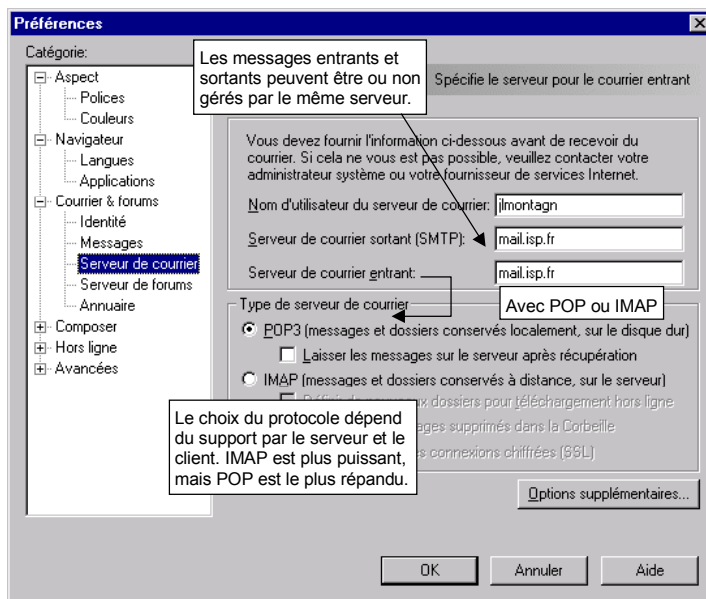
Par exemple, si le serveur SMTP de dépôt des messages de votre ISP est différent du serveur POP, vous pouvez l'indiquer en cliquant sur "Options avancées..."

De même, dans l'onglet "Connexion", vous retrouvez bien le profil défini à la section "Configurer l'accès réseau à distance" du chapitre précédent, qui indique quel modem utiliser pour se connecter à l'Internet.

Si vous utilisez Outlook Express, vous retrouvez les mêmes options, mais dans des menus différents.



Avec le navigateur Netscape, on retrouve les mêmes paramètres dans le menu “ Edition→Préférences ”.



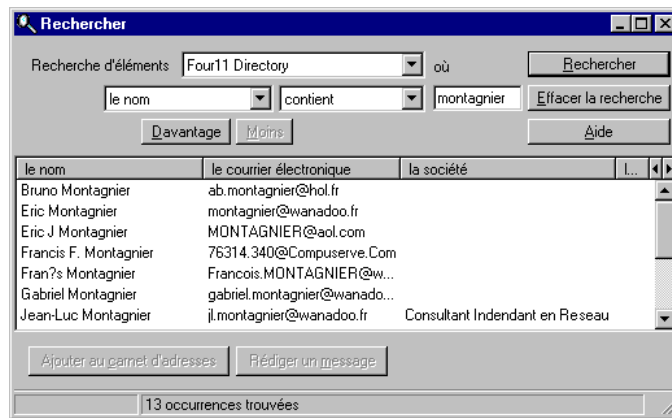
Le choix des protocoles POP ou IMAP est dicté par les possibilités offertes par votre ISP.

Les annuaires

L'Internet recèle également des serveurs d'annuaires contenant les coordonnées de la plupart des utilisateurs de l'Internet. Leur contenu est, en effet, alimenté par les informations issues de votre abonnement auprès de votre ISP. Mais vous pouvez également y trouver la plupart des abonnés au téléphone.

Les plus gros serveurs d'annuaires sont Four11 (ldap.four11.com, racheté par Yahoo!) et Bigfoot (ldap.bigfoot.com). Ils sont consultables *via* le protocole LDAP (*Lightweight Directory Access Protocol*), une version allégée de la norme X.500. Pour y accéder, cliquez sur le menu “ Edition→Rechercher dans l'annuaire ”, puis saisissez, par exemple, votre nom pour voir si vous êtes connu.

Si, lors de votre abonnement, vous avez choisi de ne pas diffuser votre adresse e-mail, vous ne figurerez dans aucun annuaire. L'avantage est que vous ne serez pas inondé de messages publicitaires (appelés *spamming* par les Américains).

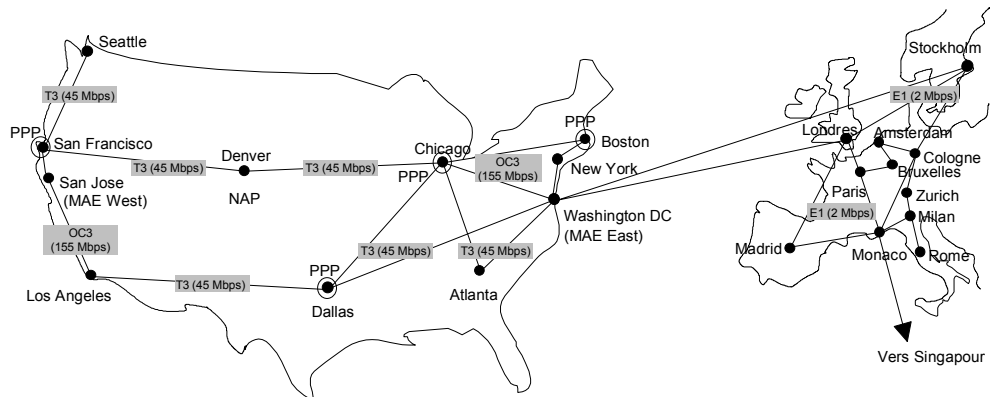


Mais qu'est-ce que l'Internet ?

L'Internet est la concaténation de différents réseaux appartenant à des opérateurs privés ou publics (par exemple, France Télécom en France). La plupart des opérateurs ne disposent pas de leur propre infrastructure (les câbles) ; ils louent tout ou partie des liaisons et y connectent leurs routeurs (la plupart du temps de marque Cisco).

Les réseaux des opérateurs se superposent, se dédoublent et se rejoignent par moments au niveau de points de concentration.

Figure 3-2.
L'Internet.



Les points de concentration reposent sur des réseaux à hauts débits (ATM, FDDI, voire Ethernet commuté) qui utilisent la fibre optique, parcourant une ville importante. Il en existe de trois types :

- Les **MAE** (*Metropolitan Area Exchange*) qui sont détenus par Worldcom (via sa filiale MFS). Il en existe sept aux États-Unis (dont les plus gros sont situés à San Jose, dans la Silicon Valley, et à Washington) et un à Paris.
- Les **NAP** (*Network Access Point*) qui sont détenus par le NSF (*National Science Foundation*) et gérés par des opérateurs privés. Il en existe actuellement quatre aux États-Unis (un à San Francisco géré par Pacific Bell, un à Chicago géré par Bellcore, un à Washington DC géré par Ameritech, et un quatrième près de New York géré par Sprint).
- Les **PPP** (*Private Peering Point*) qui sont gérés directement par des ISP qui se sont associés afin de contourner les MAE et les NAP.

Les ISP (*Internet Service Provider*) connectent leurs routeurs aux MAE, NAP et PPP (moyennant une redevance) interconnectant ainsi leurs réseaux. L'ensemble de ces interconnexions forme l'Internet.

Les liaisons entre les villes sont constituées de câbles en fibre optique ou en cuivre au bout desquels on retrouve les routeurs, équipements réseau de base qui permettent d'acheminer toutes les communications sur l'Internet. Les routeurs utilisent les protocoles Frame Relay et, de plus en plus, ATM.

Les routeurs des ISP ainsi que les commutateurs des MAE, NAP et PPP sont installés dans des locaux techniques, sans doute dans un immeuble devant lequel vous passez tous les jours sans vous en rendre compte. Certains NAP sont même perdus au fond d'un parking souterrain...

L'Internet, c'est donc cela : une collection de réseaux détenus et gérés par des opérateurs privés ou gouvernementaux (essentiellement le NSF) autour desquels gravitent les ISP qui revendent leurs services aux consommateurs que nous sommes.

La conséquence de cette situation est qu'il n'y a aucune garantie de service : les temps de réponse dépendent de la charge réseau, et les pannes sur l'Internet ne sont pas des légendes, certaines sont même célèbres :

- En 1996, une erreur de configuration sur un routeur a créé une nouvelle route redirigeant ainsi 25 % du trafic Internet sur une seule liaison à 1,5 Mbit/s. L'engorgement résultant a provoqué un arrêt quasi total de l'Internet pendant deux heures.
- En 1997, la panne d'un serveur a abîmé un fichier DNS contenant près d'un million de noms. Le technicien a ignoré l'alarme (normal, un bip de plus ou de moins...) et le fichier s'est répliqué sur d'autres serveurs DNS. Le temps de tout remettre en ordre, des dizaines de milliers de sites web n'ont plus été accessibles pendant plusieurs heures.
- Le 13 avril 1998, le backbone Frame Relay d'ATT s'est complètement arrêté suite à une erreur de configuration sur un commutateur doublée d'un bogue dudit équipement. Celui-ci a généré des messages d'alerte vers les autres commutateurs qui se sont engorgés et qui ont eux mêmes généré d'autres messages d'alerte, et ainsi de suite. En moins d'une demi-heure, le réseau était par terre. L'impact sur l'Internet a été faible, car il

existait des routes de secours (celles d'autres opérateurs !), mais les clients d'ATT ont été privés de réseau pendant 12 à 24 heures.

Cependant, depuis 1985, l'Internet n'a cessé d'évoluer afin de faire face à l'augmentation constante du nombre d'utilisateurs, et l'on peut espérer que, tant que les opérateurs gagneront de l'argent, ils assureront un minimum de qualité de service afin d'éviter de perdre leurs clients.

Quelques chiffres

L'Internet est en perpétuelle évolution, ce qui fait que les statistiques sont difficiles à obtenir. Les chiffres diffèrent selon les sources ; de plus, ils ne s'appuient pas toujours sur les mêmes critères. Même le NIC (*Network Information Center*), qui gère le plan d'adressage et de nommage, publie des informations provenant de sociétés privées qui scrutent l'Internet (dans le but de vendre leurs études de marchés). Il est vrai qu'à sa décharge, le NIC a un fonctionnement très décentralisé.

On ne peut donc que se fonder sur une photographie prise à un instant donné, et encore, assez floue. Début 1999, il y aurait ainsi eu :

- plus d'un million de domaines DNS dont 75 % dans “.com” ;
- plus d'un million et demi de serveurs web (nommés www) dont 44 % de marque Apache (un freeware) et 20 % de marque IIS (*Internet Information Server*, de Microsoft) ;
- plus de 160 000 réseaux de classe C et près de 10 000 réseaux de classe B ;
- plus de 40 millions d'ordinateurs connectés ;
- et près de 100 millions d'internautes (dont 50 aux États-Unis, 11 au Japon, 7 en Allemagne et 2 en France).

Pour vous donner une idée de la croissance phénoménale de l'Internet, il y aurait eu, début 2000, 80 millions d'ordinateurs connectés dont 6 millions d'internautes en France...

Vous pouvez prendre connaissance des statistiques les plus récentes en consultant, par exemple, le site <http://www.isc.org/ds>.

Qui gère l'Internet ?

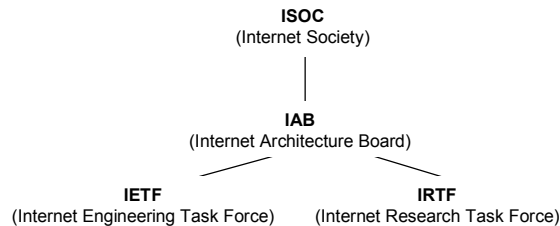
L'Internet, c'est d'abord un réseau. C'est aussi un ensemble de protocoles (plusieurs centaines !) couramment appelés protocoles Internet ou protocoles TCP/IP. Ils couvrent aussi bien les couches réseaux que les applications, telles que la messagerie ou le web. Tous ces protocoles respectent des standards définis dans des documents techniques appelés **RFC** (*Request For Comments*).

L'Internet, c'est également une communauté regroupant des organismes de recherche, des universités, des constructeurs de matériels, des éditeurs de logiciels, des opérateurs et, de plus en plus, des sociétés qui veulent simplement gagner de l'argent.

Pour organiser tout cela, l'Internet est structuré en plusieurs organisations, chacune ayant un rôle bien défini.

Figure 3-3.

Les organismes chargés de développer les protocoles Internet..



L'**ISOC** (*Internet Society*) est une association ayant pour but de promouvoir l'Internet et d'en financer le développement. Son rôle est donc de coordonner et de financer les organismes qui régulent l'Internet, tels que l'IETF. L'ISOC entretient également des relations avec l'ITU (*International Telecommunication Union*, organisation dépendant de l'ONU), les acteurs de l'industrie (constructeurs et éditeurs de matériels et de logiciels réseaux) et les gouvernements, afin d'officialiser les relations et de collecter des fonds.

Le résultat concret de ces échanges est l'apparition des RFC dans les recommandations de l'ITU-T (ITU, secteur des télécommunications) et une plus grande diffusion des protocoles édictés par cet organisme, dont les documents sont longtemps restés inaccessibles au public.

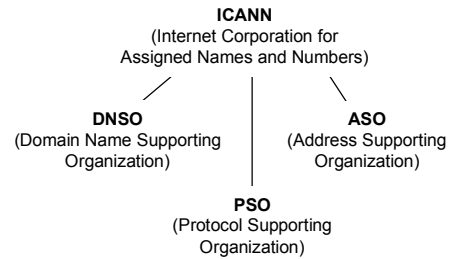
L'**IAB** (*Internet Architecture Board*) est un comité de quelques personnes qui décide des évolutions de l'Internet, telles que l'adressage, la mise en chantier d'IPv6, l'évolution de l'architecture du réseau et du DNS ou encore la sécurité. Parmi la quinzaine de membres que compte ce comité on trouve des représentants de Cisco, 3com, Microsoft, Netscape, Sun, MCI, ATT, IBM, plus quelques autres représentants d'universités américaines. La plupart d'entre eux viennent de l'IETF, les autres de l'IESG, du IANA et de l'IRTF.

L'**IETF** (*Internet Engineering Task Force*) regroupe des ingénieurs de divers horizons (instituts de recherche, universités, constructeurs et éditeurs, etc.) qui travaillent à l'élaboration des protocoles utilisés sur l'Internet. Les résultats de ses travaux aboutissent à la rédaction des **RFC** (*Request For Comments*) approuvés par l'IESG, puis validés par l'IAB, et enfin estampillés et diffusés par l'ISOC. L'IESG (*Internet Engineering Steering Group*) est le comité de validation technique de l'IETF.

L'**IRTF** (*Internet Research Task Force*) est structuré en groupes de recherche dont les objectifs sont de travailler aux protocoles de demain. L'IRTF réalise des travaux analogues à l'IETF mais sur le long terme. Le RFC 2014 décrit le fonctionnement de cette organisation. L'IRSG (*Internet Research Steering Group*) est le comité de validation technique de l'IRTF.

Figure 3-4.

Les organismes chargés de réguler l'Internet.

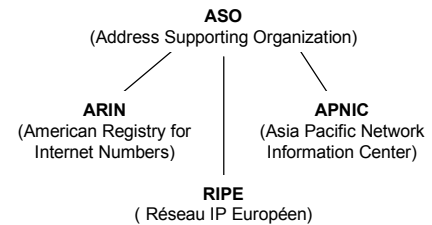


L'**ICANN** (*Internet Corporation for Assigned Names and Numbers*) a en charge la gestion des adresses IP et des noms DNS (et des serveurs racines) ainsi que l'affectation des paramètres aux protocoles IP. Cet organisme est pour cela structuré en trois SO (*Supporting Organization*) : l'**ASO**, le **DNSO** et le **PSO**.

L'**ASO** (*Address Supporting Organization*) est chargé de gérer le plan d'adressage de l'Internet. Il affecte à ce titre les adresses IP et s'appuie sur des délégations régionales.

Figure 3-5.

La gestion des adresses IP sur l'Internet.

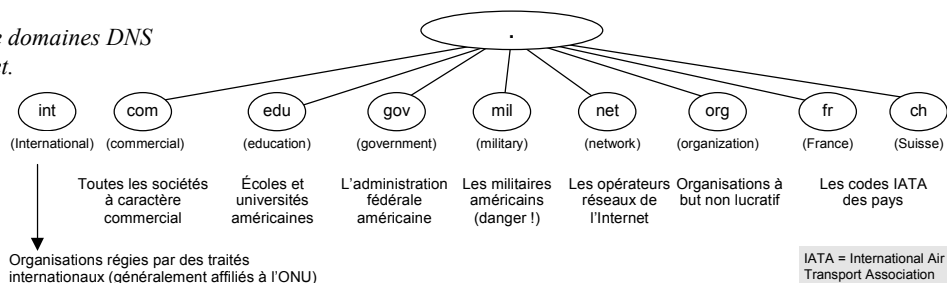


Le **DNSO** (*Domain Name Supporting Organization*) est chargé de gérer le plan de nommage de l'Internet. Il accrédite les sociétés auprès desquelles vous pouvez demander l'enregistrement d'un nom de domaine et celles qui peuvent gérer des serveurs de noms.

Les domaines de plus haut niveau sont appelés *Top Level Domains* (TLD). Parmi eux, on distingue les gTLD (*general TLD*) qui regroupent “.com”, “.edu”, etc., et les ccTLD (*country-code TLD*) qui désignent chacun un pays, par exemple “.fr”, “.uk”, etc.

Figure 3-6.

Les noms de domaines DNS sur l'Internet.



Par exemple, l'AFNIC, financé par l'INRIA et les ISP opérant en France, gère les serveurs DNS du domaine *fr* et enregistre les noms de domaines demandés par les sociétés (délai moyen : 48 heures). La liste de ces domaines est disponible à l'adresse : <ftp://ftp.nic.fr/pub/annuaire/Liste-Des-Domaines-Français>.

Le **PSO** (*Protocol Supporting Organization*) s'occupe d'enregistrer toutes sortes de valeurs utilisées par les protocoles Internet. Cela concerne, entre autres :

- les numéros des ports TCP et UDP (les *Well Known Port Numbers*) ;
- les numéros des protocoles utilisant IP (TCP=6, UDP=17, etc.) ;
- tous les codes utilisés par tous les protocoles (par exemple, les types de messages DNS, OSPF et PPP, la signification de tel octet et de tel bit, etc.) ;
- les variables des MIB ;
- les codes vendeurs des adresses MAC (repris de l'IEEE) ;
- etc.

La liste actualisée de toutes les valeurs peut être consultée sur le site web de l'ICANN, ou sur : <http://www.isi.edu/in-notes/iana/assignments/port-numbers>. La RFC 1700 regroupe par ailleurs toutes les valeurs connues en octobre 1994.

Quelques autres organismes d'intérêt général

L'**IEPG** (*Internet Engineering Planning Group*) a pour objectif de coordonner les activités des opérateurs dont les réseaux constituent l'Internet. La charte de l'IEPG fait l'objet du RFC 1690. Au sein de ce groupe, les opérateurs décident de la manière dont ils interconnectent leurs réseaux, diffusent leurs tables de routage, etc.

Dans le domaine de la sécurité, le **CERT/CC** (*Computer Emergency Response Teams/Coordination Center*) coordonne l'activité d'une dizaine de groupes de surveillance à travers le monde. Il diffuse en permanence des informations ou des bulletins d'alerte relatifs à des problèmes de sécurité (virus, bogues logiciels, attaques répertoriées, statistiques, etc.). Le **FIRST** (*Forum of Incident Response and Security Teams*) regroupe également divers groupes de surveillance, dont le CERT/CC, liés à des organismes gouvernementaux (NASA, NIH), à des universités (Oxford, Israeli Academic Network) ou à des sociétés privées.

Les anciens organismes de régulation

La réorganisation de 1999 résulte de la volonté du gouvernement américain de privatiser la gestion de l'Internet auparavant exercée par des organismes qu'il subventionnait.

L'**IANA** (*Internet Assigned Numbers Authority*) est responsable d'enregistrer toutes sortes de valeurs utilisées par les protocoles Internet. Les ccTLD étaient gérés par les délégations nationales du NIC, par exemple, l'INRIA, en France. Les domaines *gov* et *mil* sont toujours gérés par leurs organismes respectifs.

Le **NSF** (*National Science Foundation*), l'équivalent américain du CNRS, a créé le premier backbone de l'Internet dans les années 80 (il n'existe plus et a été remplacé par ceux des

opérateurs). Jusqu'à fin 98, le NSF conservait la maîtrise de l'exploitation de la partie commune de l'Internet, c'est-à-dire l'adressage IP et le nommage DNS.

Le **NIC** (*Network Information Center*) était financièrement soutenu par le NSF. Cet organisme était chargé de gérer le plan d'adressage de l'Internet, ainsi que les domaines *com*, *edu*, *net* et *org*.

Où les contacter ?

Organisme	Site web
Internet Society	www.isoc.org
IAB	www.iab.org
IETF	www.ietf.org
IRTF	www.irtf.org
Éditeur des RFC	www.rfc-editor.org
ICANN, ASO, DNSO et PSO	www.icann.org
IANA	www.iana.org
NIC en France chez les militaires américains zone Asie-Pacifique	rs.internic.net www.nic.fr ou afnic.asso.fr www.nic.mil www.apnic.net
ARIN	www.arin.net
RIPE NCC	www.ripe.net
IEPG	www.iepg.org
FIRST	www.first.org
CERT	www.cert.org

4

Construire son premier réseau local

La connexion à l'Internet a été l'occasion de se frotter aux réseaux en manipulant ses composants : modems, câbles, logiciels de communication, navigateurs et messageries.

De retour au bureau, l'étape suivante consiste à construire son propre réseau pour les besoins de son entreprise, c'est-à-dire un **intranet**. Le réseau offre, en effet, de formidables possibilités de développement : pouvoir vendre des produits au monde entier sans ouvrir de boutiques dans chaque pays, collecter des informations sur des sujets précis, échanger des documents avec ses fournisseurs, etc. Avec les réseaux, nous entrons de plein pied dans la société de l'information.

Dans une entreprise, le réseau est tout d'abord local, c'est-à-dire limité à un ou plusieurs bâtiments. Commençons donc par là.

Dans ce chapitre, vous apprendrez ainsi :

- quels sont les principes de base d'un réseau local ;
- à choisir les matériels et logiciels pour votre réseau ;
- à installer une carte réseau ;
- à configurer votre PC.

Le contexte

Avec la connexion à l'Internet vous commenciez, sans le savoir, à construire les prémices d'un réseau. Vous en avez utilisé tous les composants : câbles, matériel de connexion (dans notre cas le modem), logiciels TCP/IP, etc. Il s'agissait du type de réseau, parmi les nombreuses autres variantes possibles, qui était le plus approprié pour connecter un seul poste de travail.

Maintenant, le but est de relier plusieurs PC entre eux (de 2 à 10), par exemple, le vôtre à ceux de la secrétaire et du comptable ; ou, chez vous, entre votre bureau, la cave et la cuisine (histoire de s'amuser). Les PC sont distants de quelques mètres.

Il s'agit donc de créer un réseau adapté à ce besoin. Par conséquent, on utilisera des matériels différents de ceux du cas précédent (à chaque problème sa solution).

Le réseau qui correspond à notre situation est appelé réseau local (LAN, *Local Area Network*). Pour le mettre en place, vous avez besoin :

- d'une série de câbles qui relient les PC entre eux ;
- des cartes réseau qui permettent aux PC de se raccorder à ces câbles, d'envoyer des données et d'en recevoir ;
- de logiciels de communication, assez semblables à ceux utilisés au chapitre précédent.

QU'EST-CE QU'UN RÉSEAU LOCAL ?

Un **LAN** (*Local Area Network*) est un réseau dont la portée est limitée de quelques mètres à plusieurs centaines de mètres. C'est le type de réseau que l'on peut installer chez soi, dans des bureaux ou dans un immeuble. Un LAN, comme tout réseau, repose sur un support de transmission : un câble (en cuivre ou fibre optique) ou, plus rarement, les ondes radio.

Les réseaux locaux les plus répandus sont **Ethernet** (85 %) et **Token-Ring** (15 %). Il existe plusieurs **topologies** pour un LAN :

- En **anneau**. Les PC sont chaînés entre eux, le premier étant connecté au dernier, afin de former l'anneau.
- En **bus**. Les PC sont connectés à un câble qui parcourt tous les bureaux ou toutes les pièces de la maison.
- En **étoile**. Autour d'un équipement spécifique appelé **concentrateur** (couramment appelé *hub* pour Ethernet et *MAU* pour Token-Ring).

La topologie en étoile est la plus courante, tandis que le bus est le moyen le plus simple pour construire un réseau Ethernet.

On retrouve les mêmes briques à assembler que dans le cas précédent. La carte réseau remplace ici le modem qui servait à se connecter à l'Internet.

QUELLES DIFFÉRENCES ENTRE ETHERNET ET TOKEN RING ?

Le principe d'Ethernet repose sur un **bus partagé** : chaque station émet quand elle le souhaite mais, quand deux stations émettent en même temps, il se produit une **collision** matérialisée par la somme des deux signaux véhiculant les deux trames. Dans ce cas, les émissions sont stoppées et au bout d'un laps de temps aléatoire, une autre tentative est faite.

Le principe de Token Ring repose sur un **anneau** : chaque station attend de disposer d'un **jeton** (matérialisé par une trame d'un format particulier) avant d'émettre une trame. Le jeton circule de station en station, formant un anneau.

Le bus partagé à détection de collision et l'anneau à jeton sont deux **méthodes d'accès** à un support de transmission tel qu'un câble.

À l'inverse du bus partagé dont l'accès est **aléatoire**, la technique du jeton est **déterministe** : chaque station parle à tour de rôle au bout d'un laps de temps fixe qui dépend du nombre de stations (le temps pour le jeton de faire le tour de l'anneau). La bande passante est mieux exploitée avec Token-Ring, ce qui le rend plus performant.

L'avantage technique offert par le Token Ring n'est pas utile aux réseaux locaux. De plus, il nécessite des composants électroniques plus complexes et donc plus chers à fabriquer. En résumé, Ethernet est plus simple, plus évolutif et présente le meilleur compromis coût/performances.

Les choix de base

Quel réseau ?

Tout d'abord, quel type de réseau retenir ? Ethernet (gestion des collisions sur un bus) ou Token-Ring (gestion d'un jeton sur un anneau) ?

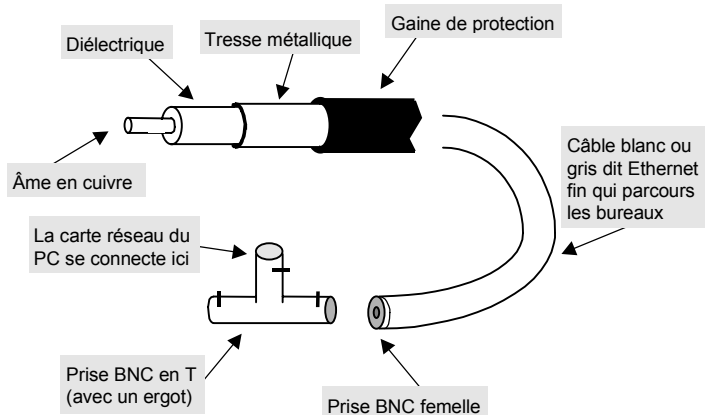
Question performances, les deux se valent, même si, à débit égal, il y a un léger avantage à utiliser Token-Ring. Cependant, Ethernet détient plus de 85 % du marché et a toujours été techniquement en avance sur Token-Ring. Si l'on doit créer soi-même un réseau à partir de rien, autant se lancer dans Ethernet : c'est plus simple et cela coûte moins cher.

Si, dans une entreprise, Token-Ring est déjà bien implanté, on peut envisager de poursuivre dans cette voie. Mais une migration vers Ethernet est toujours envisageable : tout n'est qu'une question de retour sur investissement.

Quelle topologie ?

Historiquement, le bus a été la première topologie pour Ethernet : elle repose sur un câble spécifique en cuivre, appelé câble coaxial, qui parcourt tous les bureaux dans lesquels il y a un PC à connecter.

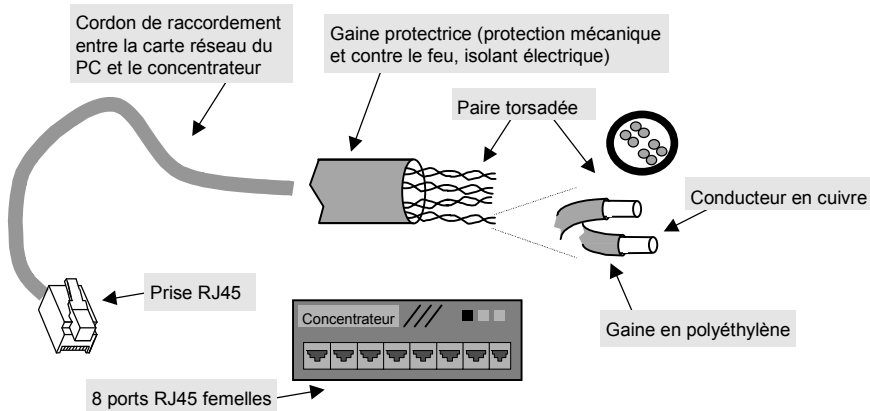
Figure 4-1.
Composant d'un réseau
Ethernet en bus.



Le câble est dit « Ethernet fin » par comparaison à une autre variante d'Ethernet, de moins en moins répandue, qui utilise un câble plus épais de couleur jaune.

Aujourd'hui, la topologie la plus répandue est celle de l'étoile qui consiste à relier tous les PC à un équipement central appelé concentrateur (*hub*, en anglais). Le câble est constitué de quatre paires de fils de cuivre torsadés et est terminé par des connecteurs RJ45.

Figure 4-2.
Composant
d'un réseau Ethernet
en étoile.



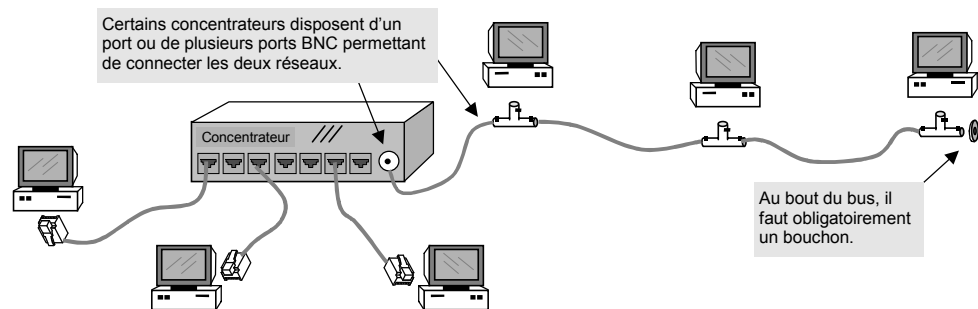
Il existe de nombreuses variantes de câbles de fils de cuivre en paires torsadées selon l'impédance, le diamètre des fils et la nature des protections. Elles seront étudiées au chapitre suivant.

Ethernet	Bus	Étoile
Câble cuivre	Coaxial	Paires torsadées
Connecteurs	BNC	RJ45
Vitesse	Limité à 10 Mbit/s	10 Mbit/s et plus
Modification du réseau	Difficile	Très facile
Remarque	De moins en moins répandu	Nécessite un concentrateur Ethernet
Adapté aux...	Petits réseaux locaux	Petits et grands réseaux locaux

En définitive, l’Ethernet en bus est la solution la plus économique lorsque l’on veut connecter quelques PC qui sont regroupés dans une seule pièce. L’Ethernet en étoile est plus cher puisqu’il nécessite un concentrateur (de 500 à 2 000 F en entrée de gamme selon le nombre de ports RJ45).

À moins que vous ne disposiez de matériel de récupération de type BNC, la topologie en étoile est conseillée. En effet, elle vous permettra de faire évoluer votre réseau tout en conservant les cartes et le concentrateur.

Figure 4-3.
*Réseau Ethernet
en étoile et en bus*



De quoi a-t-on besoin ?

De cartes réseau

Chaque PC a besoin d'un équipement capable de « parler Ethernet » : c'est le rôle de la carte réseau, dite carte Ethernet, et souvent appelée NIC (*Network Interface Card*). Elle s'insère dans un emplacement (*slot*) du PC qui lui est réservé.

Il existe plusieurs types de cartes Ethernet qui se distinguent par leur connecteur :

- BNC pour l'Ethernet fin en bus ;
- RJ45 pour l'Ethernet en étoile ;
- AUI pour l'Ethernet en bus ou en étoile.

Certaines cartes proposent une combinaison de deux de ces prises, voire les trois.

La prise AUI (*Attachment Unit Interface*) permet de connecter un équipement appelé *transceiver*, qui réalise l'adaptation au câble. Il existe ainsi des *transceivers* de types BNC, RJ45 et en fibre optique. L'acquisition de cette carte (200 F environ) peut être envisagée si votre réseau nécessite plusieurs types de câbles (coaxial, en paire torsadée, en fibre optique), voire un support de transmission radio (très peu répandu). Dans le cas de cartes RJ45 ou BNC, le *transceiver* est intégré à la carte.

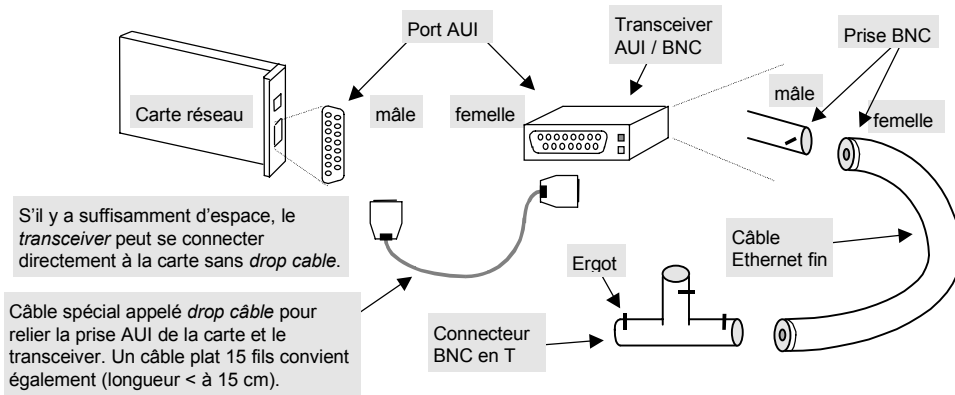
QU'EST-CE QU'UNE CARTE ETHERNET ?

L'ordinateur traite les informations sous forme numérique et sous forme de mots de 32 ou 64 bits (64 éléments d'informations binaires — 0 ou 1).

Une carte réseau Ethernet permet de convertir ces informations en signaux électriques qui sont émis sur le câble. La manière de représenter les bits d'information en signaux s'appelle le **codage**. Pour Ethernet, il s'agit du codage **Manchester**.

La carte envoie ces bits par groupes, appelés **trames Ethernet**. La norme Ethernet spécifie les couches 1 (physique : transmission des signaux par la carte réseau) et 2 (logique : format des trames Ethernet).

Figure 4-4.
Connectique AUI / BNC.



Si vous démarrez avec un réseau en bus, il est conseillé d'acheter une carte équipée de deux connecteurs, un BNC et un RJ45 (la différence de coût est minime). Cela vous permettra de la réutiliser si vous changez de réseau.

Le coût d'une carte dépend de ses performances, et notamment du bus :

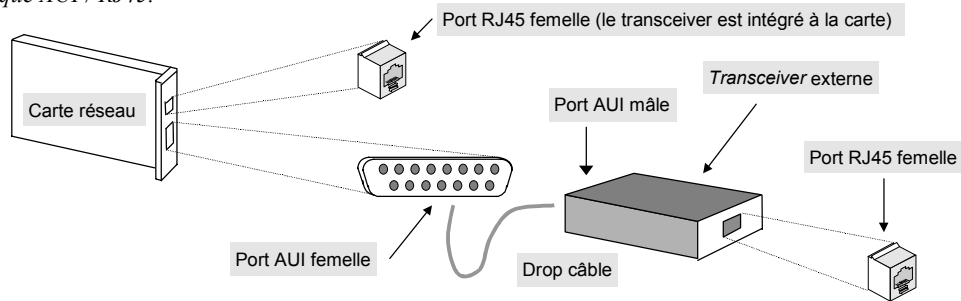
- entrée de gamme avec bus PCI : de 250 à 600 F HT ;
- entrée de gamme avec bus ISA (moins performante que PCI) : de 250 à 800 F HT.

Dans le haut de gamme, les cartes performantes sont celles qui échangent les données avec l'ordinateur *via* un DMA (*Direct Memory Access*), composant électronique spécialisé.

Pour notre réseau, une carte d'entrée de gamme suffira. Les cartes haut de gamme sont plutôt destinées aux serveurs.

Figure 4-5.

Connectique AUI / RJ45.



De cordons de raccordement

Le cordon de raccordement (également appelé cordon de brassage) est nécessaire pour connecter chaque carte réseau au concentrateur. Pour notre réseau, le plus simple est de poser un câblage volant, c'est-à-dire constitué uniquement de cordons de brassage. D'autres types de câblages — plus complexes et plus chers — seront étudiés au chapitre suivant, car au-delà de dix PC le câblage volant devient ingérable et source de problèmes.

Le support de transmission que nous avons choisi est un câble cuivre en paires torsadées dont chaque extrémité est pourvue d'une prise RJ45. Sa longueur ne doit pas excéder cent mètres, selon la qualité du câble et l'environnement électrique. Pensez notamment à éloigner vos câbles de toutes sources de perturbations : appareils électriques tels que la cafetière, le ventilateur, l'aspirateur, le moteur de l'ascenseur, le transformateur de courant, etc.

Il existe de nombreux types de câbles. Toutefois, pour notre premier réseau, le choix n'a guère d'importance. Précisons simplement, et sans entrer dans les détails, qu'il est conseillé d'acheter un câble UTP (*Unshielded Twisted Pair*), 100 Ohms, catégorie 5. C'est le moins cher, il répond à des normes précises et il est parfaitement adapté à nos besoins.

LES CÂBLES CUIVRE EN PAIRES TORSADÉES

Les câbles se différencient, avant tout, par leur **impédance**, exprimée en Ohms (Ω). Les valeurs rencontrées pour les réseaux locaux sont : 100, 120 et 150 Ohms. Plus l'impédance est élevée, meilleure est la qualité du câble (le signal est moins affaibli), mais plus son coût est élevé. Le plus répandu est le **100 Ohms**.

Les câbles se différencient également par la présence ou non de protection contre les perturbations émises par les courants électriques. Il existe des câbles sans protection, dits **UTP** (*Unshielded Twisted Pair*), avec un **écran**, dits **FTP** (*Foiled Twisted Pair*) et avec un **blindage**, dits **STP** (*Shielded Twisted Pair*). Il existe aussi la combinaison **SFTP**.

La prise la plus répandue pour les câbles en paires torsadées est la **RJ45** (*Registered Jack 45*), normalisée ISO 8877.



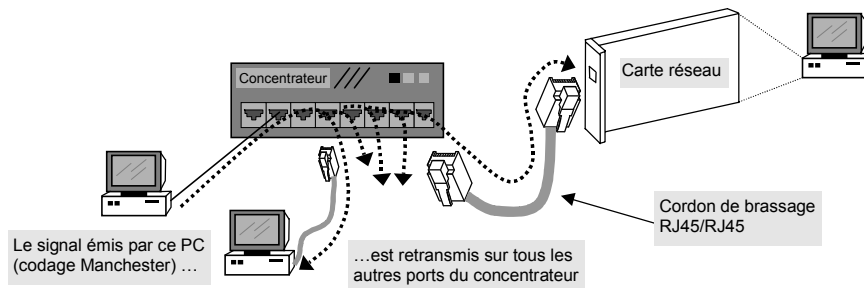
Le câble utilisé entre le PC et le concentrateur doit être droit ; les fils émission et réception ne doivent pas être croisés. Le croisement est, en effet, réalisé dans la prise RJ45 du concentrateur. Pour s'en assurer, il suffit de mettre l'un à côté de l'autre les deux connecteurs RJ45 du câble, orientés dans le même sens, et d'examiner la couleur des huit fils (généralement, le connecteur RJ45 est transparent). Si le câble est droit, les couleurs apparaissent dans cet ordre : bleu, orange, noir, rouge, vert, jaune, marron et gris. Ethernet utilise les fils numérotés 1,2 (émission) et 3,6 (réception).

Si, en revanche, vous connectez deux PC directement (sans concentrateur), le cordon doit être croisé (ce qui est logique).

D'un concentrateur

Le concentrateur est un appareil qui régénère les signaux. En effet, le signal émis par la carte Ethernet s'affaiblit en parcourant le câble et, au-delà de cent mètres, il peut devenir trop faible. Cette distance correspond en fait au maximum autorisé par la norme entre un PC et le concentrateur. Un signal émis par un PC est régénéré sur tous les autres ports du concentrateur (il joue le rôle de répéteur).

Figure 4-6.
*Fonctionnement
d'un réseau Ethernet
en étoile.*



Cela nous amène à la remarque suivante : s'il n'y a que deux PC à connecter, un concentrateur est inutile ; les deux cartes réseau peuvent être reliées directement *via* un cordon de brassage n'excédant pas cent mètres. Un câble ne disposant que de deux extrémités (loi physique incontournable de notre univers), la connexion de trois PC ou plus passe obligatoirement par un concentrateur

La plupart des concentrateurs sont dits « intelligents ». Cela signifie qu'ils disposent de mécanismes permettant de détecter les erreurs (signal trop faible, collisions, etc.) et de désactiver le port par lequel ces erreurs ont été détectées afin de ne pas perturber les autres ports (fonction de partitionnement).

Quel concentrateur choisir ? Dans notre cas, un modèle d'entrée de gamme est nettement suffisant. Le critère de choix est alors le nombre de ports RJ45, qui conditionne le nombre de PC à connecter : 4, 5, 6, 8, 12, 16, 24, 32 et 48 ports sont des valeurs couramment proposées.

En dehors du nombre de ports, les concentrateurs se distinguent par différentes fonctions.

Certains sont dits administrables ou *manageables*). Cela signifie qu'ils sont équipés d'un logiciel SNMP (*Simple Network Management Protocol*) qui permet de les administrer à distance. Dans notre cas, cette fonction n'est pas indispensable, d'autant plus que la différence de prix peut atteindre 1 000 F.

D'autres sont dits empilables (*stackable*) : cela veut dire qu'ils peuvent être chaînés afin d'augmenter le nombre total de ports. Le chaînage est effectué à l'aide d'un bus souvent matérialisé par un câble externe spécifique. Il n'existe aucune norme en la matière, ce qui signifie que vous ne pouvez pas chaîner deux concentrateurs de marque différente (un 3com avec un Dlink, par exemple).

Dans notre cas, cette fonction n'est pas intéressante, car il existe un autre moyen de chaîner les concentrateurs. Les concentrateurs sont, en effet, couramment équipés d'un port « *uplink* » de type RJ45 et/ou AUI et/ou BNC (attention au choix !) qui permet de chaîner les concentrateurs.

QU'EST CE QU'UN SEGMENT ?

Un concentrateur Ethernet (*hub*) concentre les connexions réseau des PC pour former un **segment** Ethernet.

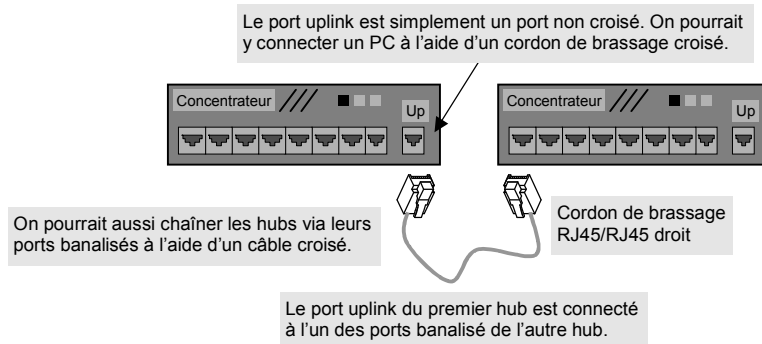
Au sein d'un segment, toutes les trames émises par un PC sont transmises par l'intermédiaire d'un *hub* à tous les autres ports (qu'un PC soit ou non connecté). Cela signifie que si deux PC émettent en même temps, la somme des deux signaux générés excédera la limite permise par la norme, ce qui correspondra à une collision. Un segment délimite donc un **domaine de collision**.

QU'EST CE QUE SNMP ?

Le protocole SNMP (*Simple Network Management Program*) permet d'interroger, de configurer et de surveiller à distance un équipement réseau. Un logiciel serveur, appelé **agent SNMP**, est implanté dans l'équipement à gérer, par exemple un concentrateur. Cet agent répond aux requêtes de clients situés dans les **stations d'administration**.

Figure 4-7.

Chaînage des concentrateurs.



Le bus externe permet à une pile de concentrateurs d'être vue comme étant un seul et unique élément, ce qui peut être utile pour l'administration à distance *via* SNMP.

Les *stackables* offrent également d'autres fonctions, telles que la segmentation port par port ou par groupes de ports. Un réseau Ethernet est constitué d'un segment matérialisé par le câble du bus ou le concentrateur de l'étoile. Un concentrateur segmentable permet de créer plusieurs segments Ethernet indépendants : un logiciel interne permet d'affecter un port (parfois cette action n'est possible que par groupes de 4 ou 8 ports) au segment Ethernet 0, et l'autre au segment Ethernet 1. Aucun trafic ne passe entre les deux réseaux ; les PC situés sur des segments différents ne peuvent donc pas communiquer entre eux.

L'intérêt de la segmentation est de créer des réseaux protégés (un pour la comptabilité séparé des autres, par exemple) ou de pallier un problème de charge : s'il y a trop de trafic sur un segment, il est possible de segmenter le réseau en répartissant les PC de part et d'autre en fonction de leur besoin de communication.

Fonctionnalité	Description	Intérêt	Coût pour 8 ports en F HT
Intelligent	Partitionnement des ports	Limite la portée d'un problème	De 500 à 1 000
Administrable	Gestion à distance <i>via</i> SNMP	Intéressant pour les grands réseaux	De 1 000 à 2 000
Empilable (+ administrable)	Chaînage <i>via</i> un bus propriétaire	Traité comme étant un seul <i>hub</i> administrable	De 2 000 à 3 000
Autres fonctions justifiant les différences de prix			
Port uplink	Chaînage <i>via</i> un port dédié	Augmentation du nombre de ports	± 200
Segmentation	Répartition de la charge sur plusieurs segments	Souplesse d'évolution	± 1 000
Un ou deux slots d'extension	Ajout de ports en fibres ou autre	Souplesse d'évolution	± 1 000

Pour notre premier réseau, un concentrateur dépourvu de fonction spécifique (c'est-à-dire non empilable, non administrable, sans segmentation, etc.) convient parfaitement. Il couvre tout à fait les besoins d'un particulier, d'une association, d'un cabinet de profession libérale, etc., c'est-à-dire tous les cas où vous êtes certain que le nombre de PC ne dépassera jamais 3, 16 ou 32 postes. Il suffit d'acheter le *hub* qui offre la bonne modularité.

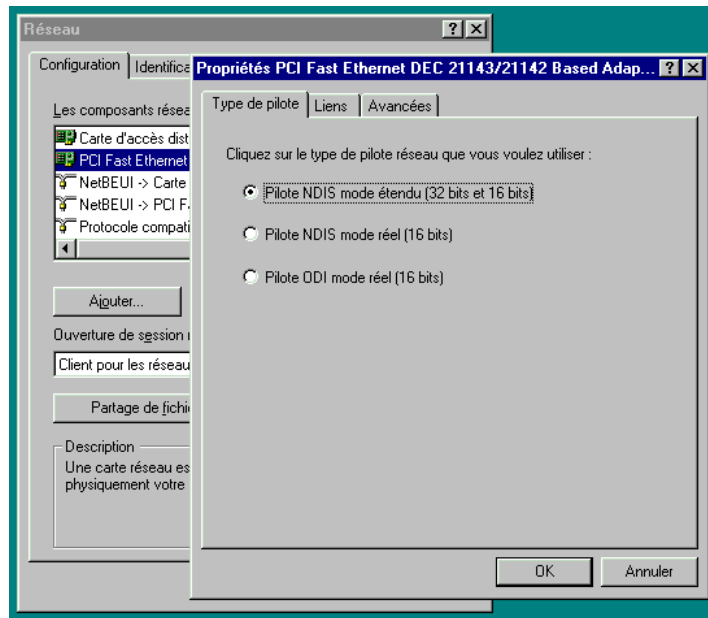
De logiciels de communications

Tout réseau nécessite du matériel et des logiciels, à l'instar d'une connexion à l'Internet qui requiert un *driver* (pour piloter la carte réseau), une pile TCP/IP et au moins un navigateur.

Le pilote est fourni par le constructeur de la carte. Cependant, si ce dernier a passé des accords avec Microsoft, il sera fourni en standard avec Windows. C'est le cas, par exemple, des cartes 3com, Dlink, HP, etc.

La pile TCP/IP est la même que celle utilisée avec l'Internet. En effet, n'importe quel type de carte peut être utilisée avec différentes piles TCP/IP du marché (celles de Windows — Netmanage — de FTP software — WRQ, etc.). Cela est rendu possible grâce à une interface d'accès standardisée sous Windows, appelée NDIS (*Network Driver Interface Specification*).

Par exemple, lors de l'installation de l'accès distant, Windows a installé un driver NDIS pour votre modem — driver fourni par le constructeur ou livré en standard avec Windows — qui dispose d'une interface NDIS. Vous pouvez le vérifier en allant dans le menu “Démarrer→Paramètres→Panneau de configuration→Réseau” qui affiche l'écran suivant (Windows 95).



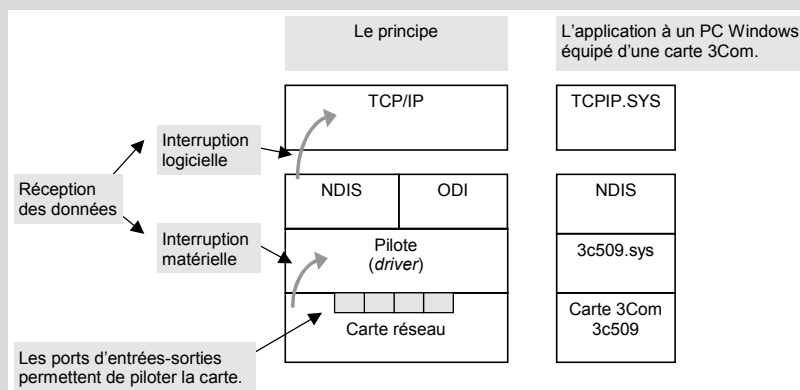
Ainsi, TCP/IP utilise-t-il les mêmes commandes NDIS, quel que soit le périphérique à piloter (une carte réseau, un modem, etc.), grâce à un driver propre à chaque matériel mais qui respecte la même interface logicielle.

LE POINT SUR LES DRIVERS

Le pilotage de chaque carte varie d'un constructeur à l'autre. Il n'existe pas, comme c'est le cas des PC, de standard relatif à la compatibilité matérielle des cartes (il n'y a pas d'« Intel Inside » !). Pour cette raison, chaque carte nécessite un pilote adapté.

À un moment ou à un autre, il faut quand même respecter un standard pour que la même pile TCP/IP puisse dialoguer avec n'importe quel pilote. La standardisation est réalisée au niveau de l'interface d'accès au pilote : la couche TCP/IP donne ainsi des ordres au pilote (du type « envoie une trame ») via une interface unique.

Dans le monde des PC, ces interfaces sont appelées **NDIS** (*Network Driver Interface Specification*) par Microsoft et **ODI** (*Open Data-link Interface*) par Novell. Ces deux standards étant bien sûr incompatibles, les cartes sont donc livrées avec deux versions du même driver : l'un avec une interface NDIS, l'autre avec une interface ODI.



Une fois l'interface d'accès au pilote standardisée, n'importe quel logiciel réseau peut être implanté. Le principe consiste à ouvrir un lien au moyen d'interruptions logicielles, de mémoires partagées, de descripteurs de tampons, etc. L'opération s'appelle **bind** (liaison), et le lien un **SAP** (*Service Access Point*).

Au-dessus de TCP/IP, on retrouve des applications diverses, tel que notre navigateur Internet qui pourra être utilisé sur notre réseau local, que l'on appellera alors intranet. D'autres applications sont possibles en local, à commencer par le partage des fichiers et l'impression. Dans l'environnement Windows, cela est réalisé par un protocole appelé Netbios (*Network Basic Input Output System*) qui fonctionne au-dessus de TCP/IP. On ne retrouve pas Netbios sur l'Internet. En effet, d'une part il est spécifique à Microsoft, d'autre part il n'est pas adapté à

ce type de réseau. Netbios est avant tout conçu pour le réseau local et devra être installé sur le PC en même temps que TCP/IP.

Comment faire fonctionner tout cela ?

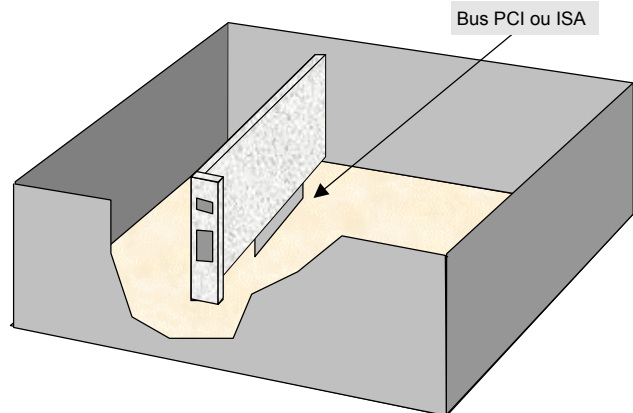
Maintenant que vous avez acheté les cartes réseau (une par PC à connecter), un concentrateur (ou plusieurs !) ainsi que les câbles, il ne reste plus qu'à assembler tout cela.

Installer les cartes réseau et les drivers

La première chose à faire est d'installer les cartes sur chaque PC. La procédure est standard, mais certaines subtilités — telles que le positionnement de cavaliers (*jumpers*) ou de commutateurs (*switches*) — sont à prendre en compte. Généralement, il n'y a rien à configurer avec Windows 9.x ; ce dernier reconnaît automatiquement la carte et la configure avec les bons paramètres. La documentation livrée avec la carte indique la procédure à suivre.

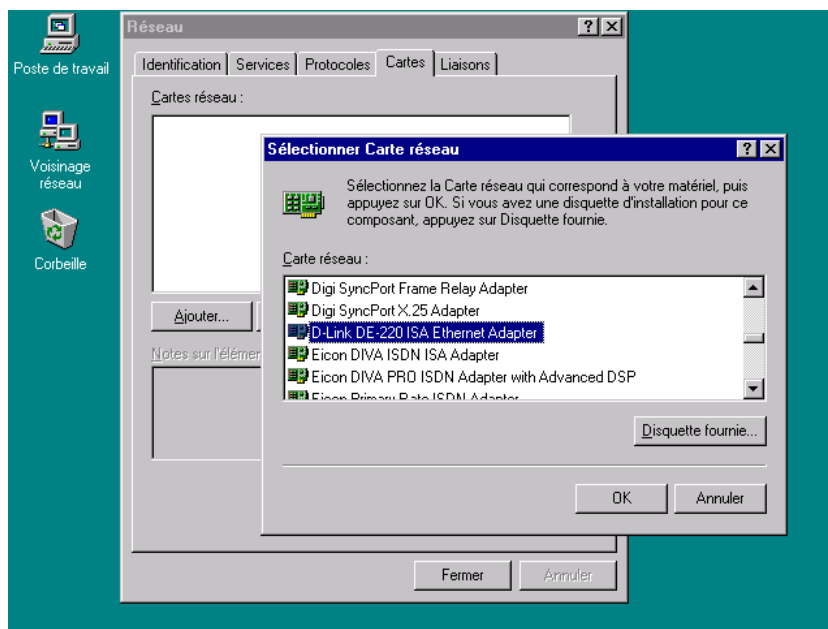
Figure 4-8.

Insertion d'une carte réseau dans un PC.



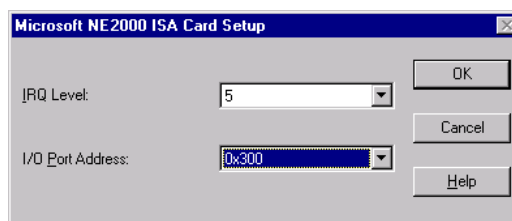
L'étape suivante consiste à installer le driver de la carte.

Si celle-ci n'est pas détectée par la fonction *Plug and Play* de Windows, vous pouvez lancer la procédure en cliquant sur le menu "Démarrer→Paramètres→Panneau de configuration→Réseau". Vous obtenez alors l'écran présenté à la figure 4-9 (quasi identique sous Windows NT et Windows 9x).



Le programme propose de choisir le driver dans une liste. S'il n'y figure pas, cliquez sur "Disquette fournie...". Cliquez ensuite sur "OK".

Selon les cartes, une boîte de dialogue vous demande de spécifier la valeur de certains paramètres.



Ces paramètres concernent :

- les interruptions matérielles, appelées IRQ (*Interruption Request*), utilisées par la carte pour avertir le driver qu'une trame vient d'arriver, par exemple ;
- les ports d'entrée-sortie (*I/O port, Input Output*) qui correspondent à des registres (mémoire partagée par la carte et le PC) permettant au driver d'envoyer des commandes à la carte (envoyer une trame, par exemple) ;
- le port DMA (*Direct Memory Access*) si la carte utilise ce mode.

La documentation vous indique la marche à suivre. Généralement, les valeurs par défaut conviennent. Elles doivent être modifiées seulement si vous possédez d'autres cartes qui utilisent les mêmes IRQ et/ou ports I/O.

Cliquez sur "OK" pour terminer l'opération. Le PC affiche alors une série d'écrans et vous demande de réinitialiser l'ordinateur.

Configurer les adresses IP

À la différence de la connexion Internet, il n'existe pas d'ISP pour attribuer automatiquement des adresses IP. À présent, vous êtes chez vous, sur votre réseau, et vous êtes seul maître à bord.

Il faut donc affecter vous-même une adresse IP à chaque poste de travail afin qu'il puisse être identifié de manière unique.

À QUOI SERT L'ADRESSAGE ?

Comme pour le courrier postal, l'adresse permet d'acheminer les trames Ethernet et les paquets IP.

Les réseaux Ethernet utilisent un **adressage plat** : les cartes réseau sont identifiées par une adresse unique, l'adresse MAC (de niveau 2).

Le protocole IP utilise, quant à lui, un **adressage hiérarchique** (de niveau 3) structuré en un **numéro de réseau** et un **numéro de station** au sein de ce réseau (32 bits en tout). L'adresse IP est indépendante de l'adresse MAC : un segment Ethernet peut comprendre plusieurs réseaux IP, et inversement.

Aussi bien au niveau MAC que IP, il existe trois types d'adresses :

- l'adresse **unicast** qui est affectée à une station ;
- l'adresse **multicast** qui désigne un groupe de stations ;
- l'adresse de **broadcast** qui désigne toutes les stations sur un réseau.

Une station est configurée avec une adresse MAC (celle de la carte réseau) et une adresse IP (celle de la pile IP). Des mécanismes spécifiques permettent de réaliser automatiquement la correspondance entre les deux types d'adresses.

Chaque trame Ethernet contient l'adresse MAC de l'émetteur et celle du destinataire. De même, chaque paquet contient les adresses IP de l'émetteur et du destinataire, ce qui permet de les acheminer indépendamment les uns des autres.

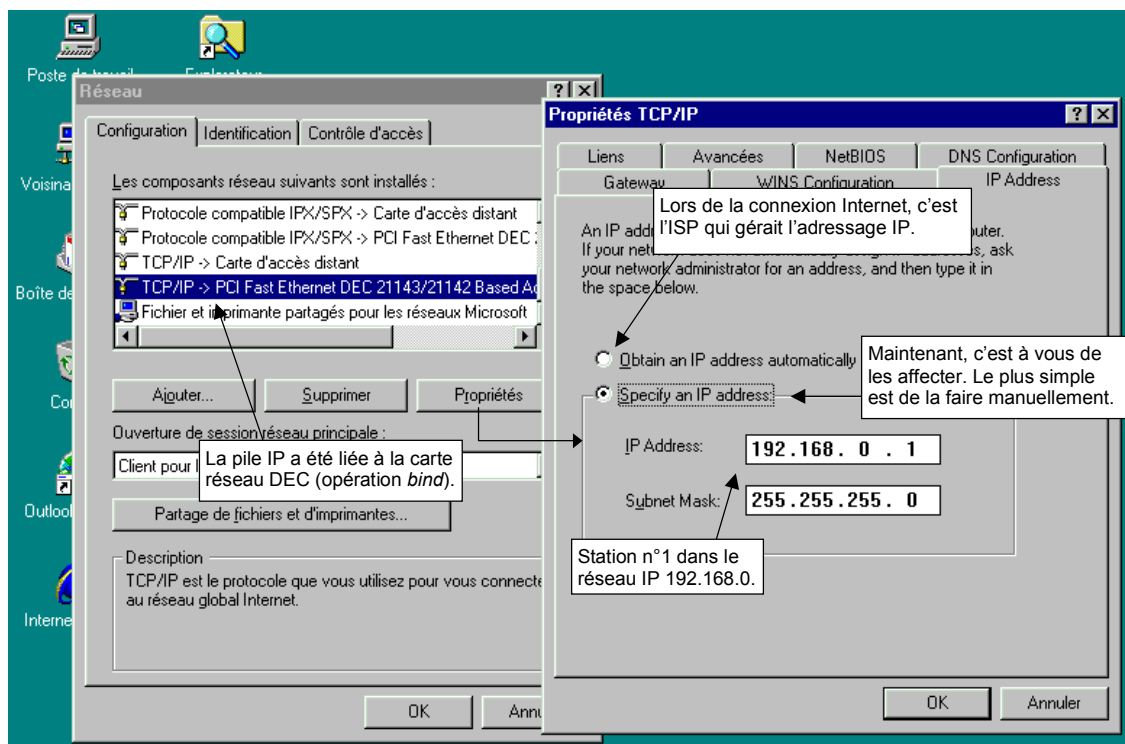
Pour reprendre l'analogie avec les adresses postales, une adresse IP est composée d'un numéro de réseau (le nom d'une rue) et d'un numéro de station au sein de ce réseau (le numéro de votre maison).

Par convention, l'adresse IP s'écrit avec quatre numéros, de 1 à 255, séparés par des points, par exemple 192.162.0.1. Une partie de cette adresse désigne un réseau, l'autre le numéro de

station au sein de ce réseau. Le protocole IP utilise un masque pour distinguer les deux parties. Dans cet exemple, il sera égal à 255.255.255.0, indiquant que les trois premiers chiffres de l'adresse désignent le numéro de réseau, et le dernier celui de la station.

Dans notre cas, il faut s'arranger pour configurer toutes nos stations dans le même réseau logique IP. Nous choisirons donc le réseau 192.168.0 et affecterons à nos PC les numéros compris entre 1 et 254, ce qui donne une plage d'adresses comprise entre 192.168.0.1 et 192.168.0.254.

Sur chaque PC (Windows 9.x), il faut donc aller dans le menu "Démarrer→Panneau de configuration→Réseau" pour configurer ces adresses. Vous obtenez alors l'écran illustré sur la figure ci-après.



Pour l'instant nous n'avons pas besoin d'en savoir plus, car nous avons créé un petit réseau. Le chapitre 7 présente, dans le détail, tous les mécanismes de l'adressage.

Installer les concentrateurs et y raccorder les PC

L'installation des concentrateurs est simple, puisqu'il n'y a aucun paramètre à configurer, ni logiciel à installer. Il suffit de les brancher sur une prise électrique et d'appuyer sur l'interrupteur. Si rien ne se produit, le matériel est en panne ☹.

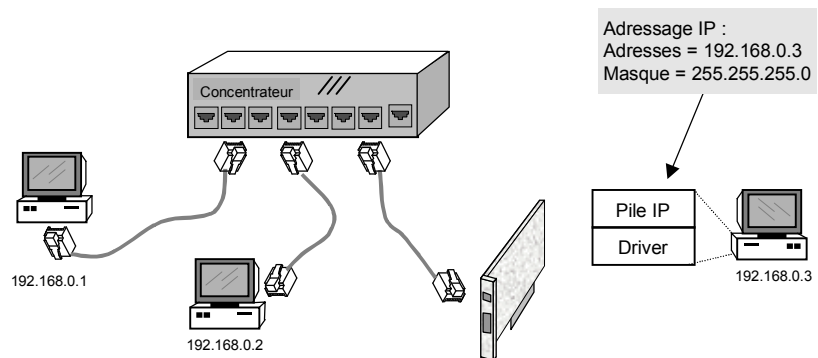
Quelques précautions doivent cependant être prises :

- Utilisez une prise électrique protégée par un disjoncteur dédié aux équipements informatiques afin d'éviter tout parasite provenant d'un autre appareil électrique (cafetière, aspirateur, etc.).
- Placez le concentrateur en hauteur, dans un endroit aéré et éloigné de toute source électrique importante (moteur d'ascenseur, encore la cafetière, etc.).

S'il est administrable, le concentrateur pourra être configuré ultérieurement pour des fonctions spécifiques liées à l'administration SNMP et à la création de segments.

La connexion des PC est également simple : il suffit de raccorder un cordon de brassage au PC et de choisir, au hasard, un des ports du concentrateur.

Figure 4-9.
*Un réseau local
Ethernet simple.*



Comme vous le voyez, ce type d'installation convient à un faible nombre de PC, de préférence regroupés dans un bureau. Rapidement, il devient nécessaire d'organiser le câblage et la mise en place des concentrateurs d'une autre manière. C'est ce que nous allons voir au chapitre suivant.

Mettre en place un système de câblage

Le câblage volant, tel qu'il a été installé pour notre petit réseau lors du chapitre précédent, ne peut pas être généralisé à grande échelle. En effet, au delà de dix postes, il devient rapidement source de problèmes : posé à même le sol, il est encombrant, voire gênant (on se prend les pieds dedans). Il est de ce fait soumis à une usure plus rapide.

Une première amélioration consiste à le faire circuler dans des parties protégées de l'immeuble : on peut le poser sous un faux plafond ou sous un faux plancher ou encore le faire passer dans une goulotte le long des murs.

Mais, à chaque nouvelle connexion ou à chaque déménagement de PC, il faut déplacer le câble et trouver un nouveau cheminement, ce qui présente des inconvénients majeurs :

- Cette opération est extrêmement difficile, voire impossible si la longueur des câbles est de plusieurs dizaines de mètres.
- Le problème d'usure demeure lorsque les câbles sont déplacés.
- Le cheminement des câbles est difficile à maîtriser : on arrive inévitablement à des situations dans lesquelles les câbles informatiques s'entrecroisent avec les câbles électriques qui sont sources de perturbations importantes. Le réseau peut ne plus fonctionner à cause de cela.

Il est donc impératif de mettre en place un système de câblage permanent (fixe et stable dans le temps) et évolutif (qui s'adapte à tous les besoins présents et futurs). Pour cela, il convient de respecter un certain nombre de règles.

Quelle est la démarche à suivre ?

Le câblage d'un immeuble requiert un certain nombre d'étapes importantes et étalées le temps :

- Si l'immeuble existe, un **audit** préalable est nécessaire afin de repérer les locaux, les sources de courants forts, les câbles existants, les cheminements possibles des futurs câbles, etc.
- La phase d'étude et d'expression des besoins, généralement appelée **APS** (avant-projet sommaire), a pour but de déterminer les spécifications fonctionnelles de l'infrastructure (locaux, gaines techniques) et du système de câblage (cuivre, fibre optique, connectique).
- Pour les grandes réalisations, l'APS n'est que l'ébauche de plusieurs scénarios. L'**APD** (avant-projet détaillé) permet alors de choisir la solution en fonction de critères techniques organisationnels et économiques.
- Le dossier de consultation (le cahier des charges) peut être formé de trois documents principaux : le **CCTP** (cahier des clauses techniques particulières), puis éventuellement le **CCTG** (cahier des clauses techniques générales) et, si vous travaillez avec l'administration française, le **CCAP** (cahier des clauses administratives particulières). Cette phase se termine par la sélection d'une entreprise de câblage.
- La phase de **suivi de chantier** nécessite un contrôle régulier et des réunions de coordination.
- La phase de réception (**recette**) consiste à tester et à valider les travaux effectués.

La première tâche est avant tout de repérer les lieux, ou de se contenter d'examiner les plans si l'immeuble n'existe pas encore.

Dans les deux cas, l'objectif est de mettre en place un câblage **systématique**, c'est-à-dire d'équiper entièrement l'immeuble. Si seuls quelques étages sont concernés, la démarche est plus ou moins la même.

Il ne s'agit donc pas de savoir où sera situé tel ou tel utilisateur, mais d'installer des prises partout dans le but de connecter n'importe qui à n'importe quelle prise pour n'importe quel type d'application. On parlera alors d'un **précâblage** multimédia ou **VDI** (voix, données, image).

L'avant-projet

Lors d'une opération de précâblage, il est important de systématiser l'implantation des prises dans tout l'immeuble. Une fois le chantier achevé, tout aménagement complémentaire sera plus délicat, plus long et plus coûteux. Le chantier de câblage est l'occasion unique de réaliser une fois pour toutes une infrastructure sans avoir à y revenir avant dix ou quinze ans.

La densité communément admise est d'environ un boîtier VDI pour 7 à 10 m² de bureaux, un boîtier pouvant regrouper de deux à quatre prises. Cette densité peut être plus élevée pour

certaines applications spécifiques comme les salles de marché : on peut trouver jusqu'à dix prises par position (occupant 3 m² environ).

On peut prendre comme repère une travée délimitée par une largeur de fenêtre. Selon les besoins, on pourra installer un boîtier VDI de deux à quatre prises par travées. Généralement, il faut une prise pour le poste de travail informatique, une autre pour le téléphone, et une troisième pour un besoin particulier (une ligne téléphonique directe, une imprimante en réseau, etc.).

Il est important de noter que le boîtier VDI doit se trouver à proximité d'un bloc de prises électriques : cela paraît une évidence, mais il faut penser à se coordonner avec l'entreprise qui réalise les travaux courants forts.

Les locaux concernés sont non seulement les bureaux, mais aussi les locaux collectifs : local photocopieur, cafétéria (on y pose souvent des bornes d'information), salles de réunions, salles de conférences, halls d'entrée (pour les bureaux d'accueil, les locaux des gardiens, etc.).

En outre, il faut aussi prévoir le câblage pour la GTB (gestion technique du bâtiment), bien que celui-ci soit souvent réalisé par une entreprise spécialisée avec laquelle il faudra de toute façon se coordonner. La GTB regroupe des besoins comme la détection incendie, les alarmes, la sécurité d'accès aux locaux, la surveillance, etc.

De même, il y a toute une série d'équipements annexes qui peuvent requérir l'emploi d'une prise :

- les téléphones d'ascenseurs ;
- les lignes directes (celles qui ne passent pas par le PABX) ;
- les bornes de réseau sans fil et de téléphonie sans fil (DECT) ;
- les badgeuses ou pointeuses ;
- etc.

Une fois l'implantation des prises définie, il faut prévoir de la place pour le cheminement des câbles et la création des locaux techniques. Le principe retenu est quasi systématiquement une topologie en étoile : les câbles relient les prises VDI à d'autres prises en local technique. Les différentes normes définissent une longueur maximale de quatre-vingt-dix mètres pour les câbles en cuivre.

De ce fait, il faut prévoir plusieurs locaux techniques au sein de l'immeuble et donc des câbles pour les relier entre eux. Plusieurs facteurs déterminent le nombre et la position des locaux techniques :

- La distance maximale de quatre-vingt-dix mètres.
- La densité des prises : on admet qu'un local peut centraliser jusqu'à 250-350 prises ; au-delà, son exploitation devient complexe (trop de câbles, trop grande concentration d'équipements).
- L'architecture des réseaux informatiques et téléphoniques : de nos jours, ils reposent sur une topologie en étoile avec des équipements installés à chaque étage et d'autres qui ont une fonction fédératrice.

L'INFRASTRUCTURE NÉCESSAIRE À UN SYSTÈME DE CÂBLAGE

Au sein d'un immeuble, de l'espace doit être réservé pour accueillir le système de câblage. Il s'agit essentiellement de **locaux techniques** et de cheminements utilisés pour relier les locaux entre eux.

Les câbles qui relient les prises VDI aux locaux techniques sont appelés **câbles de distribution**. Ceux qui relient les locaux techniques entre eux sont appelés **câbles de rocade**.

Dans les zones de circulation (couloirs, halls d'entrée, etc.), les câbles sont installés dans des **chemins de câbles** métalliques qui servent de support et offrent une protection mécanique et électromagnétique. Dans les bureaux, ces mêmes câbles sont installés dans des **goulottes** ou des tubes noyés dans le béton.

Les câbles de distribution sont généralement horizontaux et cheminent sous les **faux plafonds** et/ou sous les **faux planchers**. Ces derniers ont une fonction essentiellement esthétique et sont constitués de dalles amovibles destinées à en faciliter l'accès.

Les rocades sont verticales ou horizontales, et cheminent sous faux plafonds, faux planchers et dans des **gaines techniques** (conduits réservés aux câbles et tuyaux de toute nature).

L'architecture d'un système de câblage suit donc celle des réseaux : on définit ainsi deux niveaux de locaux techniques :

- Les **LTE** (locaux techniques d'étages) qui concentrent les prises VDI et accueillent les équipements de communication de distribution (concentrateurs, commutateurs, etc.).
- Les **LN** (locaux nodaux) qui sont reliés à tous les locaux techniques et accueillent les équipements de communication fédérateurs (PABX, commutateurs fédérateurs, routeurs, etc.).

Un troisième niveau de concentration est parfois nécessaire dans le cas où les distances sur un étage excèdent quatre-vingt-dix mètres. On peut alors trouver la dénomination de LTR (local technique rapproché) ou de LTP (local technique de proximité). Ce type d'architecture n'est cependant pas conseillée car trop complexe et mal adaptée aux architectures réseaux.

Tout ces locaux sont autant d'espaces prélevés sur la superficie utile de l'immeuble. Préparez-vous donc à quelques négociations avec l'architecte (si l'immeuble est à construire) ou avec le responsable des services généraux (s'il existe déjà). Le tableau suivant donne une idée de la surface à réserver à ces locaux techniques.

Local	Superficie moyenne
Local technique d'étage	6 m ² (3 m x 2 m)
Local nodal	24 m ² (6 m x 4 m)
Local énergie	8 m ² (4 m x 2 m)
Local opérateur	9 m ² (3 m x 3 m)
Salle informatique	28 m ² (7 m x 4 m)

Bien entendu, ces superficies doivent être ajustées en fonction du nombre de prises à câbler.

LES DIFFÉRENTS TYPES DE LOCAUX TECHNIQUES

Les LTE (**locaux techniques d'étages**) accueillent une à deux baies de câblage (distribution d'étage et rocades) ainsi qu'une à deux baies de communication (équipements de distribution des réseaux téléphonique et informatique). Généralement, un ou deux LTE par étage sont suffisants. Pour des raisons de simplicité, il faut s'arranger pour que les LTE soient tous à l'aplomb les uns des autres.

Le LN (**local nodal**) accueille les baies de câblage (distribution des serveurs et rocades) ainsi que les baies de communication (équipements centraux pour les réseaux téléphonique et informatique). Généralement, il y a deux locaux nodaux dans le bâtiment afin d'offrir une redondance pour le cheminement des câbles. Chaque LTE peut ainsi être relié aux LN via deux chemins de câbles différents. Il en est de même entre deux LTE d'un même étage.

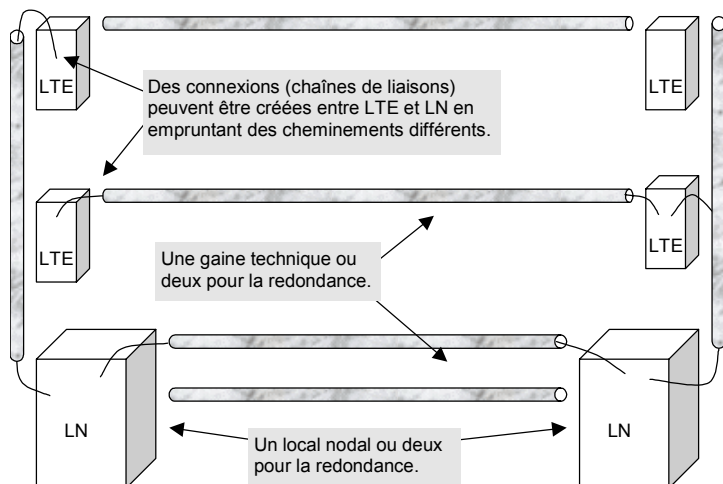
Il est généralement souhaitable de créer un LO (**local opérateur**) réservé aux arrivées télécom des opérateurs afin d'assurer une séparation claire des responsabilités entre lui et le client. Deux LO sont préférables, car la plupart des opérateurs peuvent offrir deux accès physiquement séparés et redondants. Ils doivent jouxter le ou les LN qui hébergent les équipements de communication du client (routeurs, par exemple).

Il faut aussi prévoir un LE (**local énergie**) pour accueillir l'alimentation du PABX (armoires 48v et batteries) ainsi qu'un régulateur de courant/onduleur. Il doit jouxter un local nodal. Là encore, deux locaux énergie offrent un bon niveau de redondance renforcé par deux accès EDF physiquement différents.

Une SI (**salle informatique**) accueille une à deux baies de câblage (distribution) ainsi que des serveurs informatiques. Elle doit de préférence être dédiée afin de mieux contrôler l'accès aux locaux et de séparer les responsabilités entre les équipes système et réseau. De la même manière, deux salles informatiques permettent de limiter les dégâts en cas de sinistre, et de répartir les serveurs en cluster.

Pour les petits sites, il est plus économique et plus simple de regrouper les fonctions de LN, LO et SI au sein d'un même local.

Figure 5-1.
Gaines et locaux techniques pour un système de câblage.



LES COMPOSANTS D'UN SYSTÈME DE CÂBLAGE

Les parties visibles d'un système de câblage sont les **prises utilisateur**, également appelées prises VDI (voix, données, images), installées dans les bureaux. Elles sont regroupées par blocs de 2 à 4, appelés **boîtiers VDI**. Une densité courante est d'un bloc VDI pour 7 à 10 m² de bureau.

Les prises utilisateur sont reliées en étoile à un local technique par l'intermédiaire d'un **câble** (en cuivre ou en fibre optique). Le local technique concentre 100 à 350 **câbles de distribution**, chacun se terminant par une prise identique à celle installée du côté utilisateur. Ces **prises de distribution** sont regroupées dans des **panneaux de brassage** fixés dans des **baies**.

Les prises sont reliées aux équipements informatiques et téléphoniques par l'intermédiaire de **cordons de brassage** de même nature que les câbles.

Les prises, câbles, cordons et panneaux de brassage doivent tous être issus du même constructeur afin de bénéficier de sa garantie (généralement dix à quinze ans).

Ensuite, il faut évaluer la puissance électrique consommée par les équipements informatiques. On pourra même prévoir des disjoncteurs séparés, un par baie ou pour un groupe d'équipements.

Il faut enfin prévoir une climatisation dans chaque local technique, et donc évaluer la dissipation calorifique des équipements (exprimée en Watts ou en BTU – *British Thermal Unit*). Ces valeurs sont données par les constructeurs de tout équipement informatique.

Tous ces besoins seront regroupés dans un document appelé **APS** (avant-projet sommaire) et communiqués aux corps d'état concernés (architecte, électricien, société de climatisation, etc.).

L'étude d'ingénierie

La phase d'expression des besoins est suivie d'une étude permettant d'arrêter un certain nombre de choix importants :

- Quel type de câble utiliser ?
- Quel type de prise choisir en bureau ? En local technique ?
- Où faire passer les câbles ? Comment placer les prises ?
- Où positionner les locaux techniques ? Comment les aménager ?

Quel type de câble ?

Éternel débat que celui du choix des câbles, chaque constructeur ayant des arguments en faveur de son produit. De nombreuses combinaisons techniques viennent compliquer le choix.

Pour résoudre ce dilemme, un certain nombre de questions sont à se poser, et dans le bon ordre.

Cuivre ou fibre optique ?

Les réseaux Ethernet fonctionnent sur cuivre à 10 Mbit/s et à 1 gigabit. L'avantage de la fibre optique est qu'elle permet de s'affranchir des contraintes de distance (plusieurs centaines de mètres au minimum contre quatre-vingt-dix mètres pour le cuivre). Cela tient à l'atténuation du signal, beaucoup plus important sur un câble en cuivre.

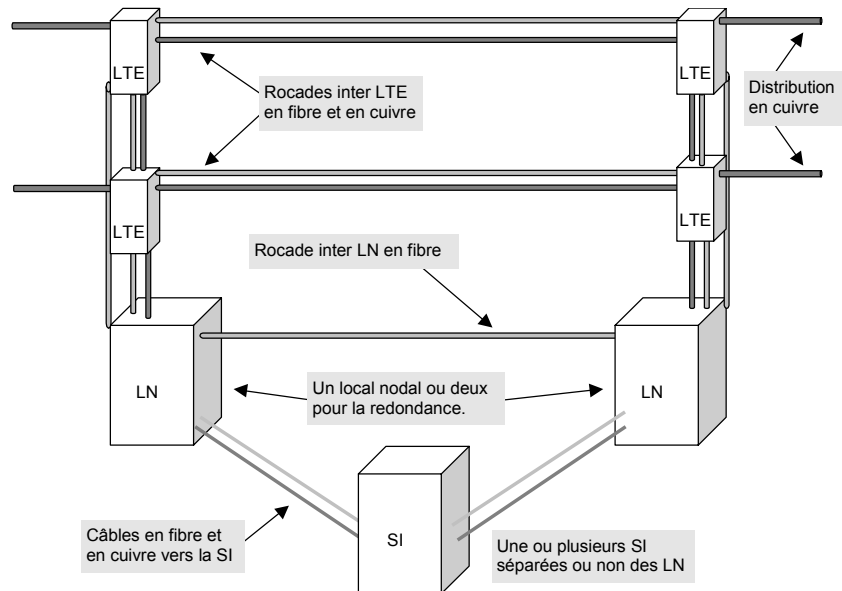
En revanche, le coût global d'un système de câblage en fibre optique est plus élevé que l'équivalent en cuivre. En effet, l'ingénierie nécessaire pour poser des câbles optiques (raccordement des connecteurs et tests) est plus complexe et plus coûteuse qu'avec des câbles en cuivre. De plus, les composants tels que les connecteurs SC et les tiroirs optiques sont également beaucoup plus chers que les prises RJ45 et les panneaux de brassage.

À titre d'indication, un système de câblage en fibre optique coûte en moyenne 60 % plus cher que l'équivalent en câble de cuivre SFTP catégorie 5E.

Il faut ajouter à cela le coût des équipements actifs (les commutateurs et cartes Ethernet), deux fois plus chers en version fibre optique, et pour une densité de ports deux fois moins élevée que leur équivalent en cuivre.

En conclusion, le câble cuivre sera privilégié pour la distribution, et la fibre optique pour la connexion entre les locaux techniques. Cette répartition des rôles offre, en outre, plus de souplesse pour positionner les LTE qui doivent être à moins de quatre-vingt-dix mètres de toutes les prises qu'ils irriguent.

Figure 2-2.
*Architecture
de câblage type.*



Pour ajouter plus de sécurité, on peut envisager de doubler les liaisons entre les LTE d'un même étage ainsi qu'entre les LN, chacune d'entre elles passant alors par deux gaines techniques différentes.

Attention, cependant, les installations téléphoniques classiques requièrent encore des connexions en cuivre entre les postes et le PABX central. Les contraintes de distance étant moins fortes (quelques centaines de mètres), il faut donc envisager du câble en cuivre, dit multipaire, entre les LTE et les locaux nodaux.

De nos jours, on privilégiera une architecture téléphonique identique à celle du réseau local avec des unités déportées dans chaque étage (de type Voice Hub, tels que proposés par Alcatel) et raccordées en fibre optique à un petit PABX central. Si vous optez pour cette solution, vous n'avez plus besoin de câbles multipaires entre le local nodal et les LTE.

Coaxial ou paires torsadées ?

Nous l'avons vu aux chapitres précédents, le câble coaxial (50 et 75 Ohms) n'est plus utilisé pour les réseaux locaux. Il pourra cependant être posé pour les besoins spécifiques de la vidéo (voir plus loin).

En revanche, la paire torsadée est le standard pour l'informatique et la téléphonie ; elle peut également être utilisée pour la distribution vidéo. Tous ces équipements (concentrateurs, commutateurs, PABX, etc.) sont, en effet, équipés de prises RJ45.

Le choix de la paire torsadée en distribution

Quelle impédance : 100, 120 ou 150 Ohms ?

Le 150 Ohms n'est qu'un artefact des réseaux Token-Ring IBM qui ne s'est jamais imposé car trop coûteux. Le 120 Ohms se voulait un compromis entre coût et performances entre les 100 et le 150 Ohms, mais ne s'est imposé qu'en France. Le 100 Ohms est le plus répandu, car il est moins cher et est soutenu par les Américains, ATT en tête.

De plus, tous les équipements informatiques sont américains et donc pourvus de connecteurs RJ45 de 100 Ohms. Cependant, un câble 120 Ohms peut y être connecté sans problème, l'affaiblissement résultant de l'adaptation d'impédance étant largement compensée par les meilleures performances du câble 120 Ohms.

En conclusion, les 100 et 120 Ohms conviennent tous deux, avec un avantage pour le premier qui est meilleur marché.

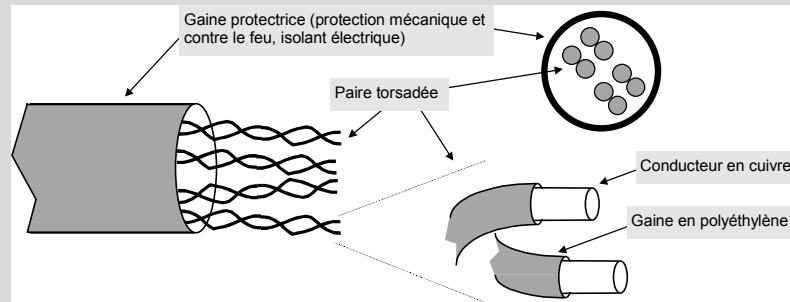
Écranté ou non ? Blindé ou non ?

Un câble UTP est non écranté et non blindé. En conséquence, l'UTP 100 Ohms est le type de câble le moins cher et donc le plus répandu.

Indiquons cependant que le coût du câble ne représente qu'une faible part (environ 10 %) du coût total d'un chantier de câblage, le plus gros morceau étant celui de la main d'œuvre pour la pose. Mais, les câbles écrantés et/ou blindés nécessitent la mise en place d'une terre informatique qui grève le budget de petites réalisations.

LE POINT SUR LA PAIRE TORSADÉE (EIA/TIA-568 ET EN50173)

Ce type de câble en cuivre comporte huit fils appariés par deux. Les deux fils de chaque paire sont torsadés selon un **pas de torsade** précis, de même que les quatre paires au sein du câble. Le but de cet arrangement est de limiter les interférences produites par chaque fil sur les autres.



La qualité d'un câble cuivre dépend bien sûr de la qualité du matériau, mais aussi des éléments suivants :

- du **diamètre** des fils, exprimé en AWG (*American Wire Gauge*) :
22 AWG = 0,63 mm ; 24 AWG = 0,5 mm (le plus courant) ; 26 AWG = 0,4 mm ;
- de l'**impédance** caractéristique, exprimée en Ohm, qui représente la résistance du câble, le 100 Ω étant la plus courante (les 120 et 150 Ω peuvent encore être rencontrés) ;
- de sa **protection** contre les champs électromagnétiques : sans protection (**UTP**, *Unshielded Twisted Pair*), avec un écran (**FTP**, *Foiled Twisted Pair*), avec un blindage (**STP**, *Shielded Twisted Pair*) ou avec les deux (**SFTP**, *Shielded & Foiled Twisted Pair*).

Les normes **EIA/TIA-568** précisent les paramètres à mesurer pour vérifier la qualité des câbles (TSB 36) et des prises RJ45 (TSB40) :

- l'**affaiblissement** du signal, exprimé en décibels pour cent mètres ;
- la **paradiaphonie NEXT** (*Near End Cross Talk*), exprimée en décibels, qui mesure la quantité de signal engendré sur une paire par une autre ;
- la **paradiaphonie cumulée** (*Powersum Next*), exprimée en décibels, qui mesure la quantité de signal engendré par toutes les paires sur une autre ;
- le **rapport signal/bruit**, exprimé en décibels.

La meilleure qualité est obtenue pour un faible affaiblissement et pour une paradiaphonie et un rapport signal/bruit élevés.

Quatre **catégories** de câbles sont définies, **5**, **5E**, **6** et **7** (les catégories 1 à 4 ne sont pas utilisées en informatique) selon la **fréquence** maximale du signal pouvant être véhiculé : cat 5 et 5E à 100 MHz, cat 6 à 250 MHz et cat 7 à 600 MHz. Plus la fréquence est élevée, plus le débit du réseau le sera. Par exemple, le 100bT fonctionne à 62,5 MHz sur deux paires, et le gigabit à 100 MHz sur quatre paires. Les normes précisent les valeurs minimale ou maximale des paramètres pour différentes fréquences de fonctionnement.

La norme européenne **EN50173** reprend le même principe, mais définit les valeurs pour une chaîne de liaison comprenant un câble de 90 mètres et deux cordons de brassages de 5 mètres chacun. On parle alors de **classes D, E, F et G**.

La compatibilité électromagnétique (EMC) — norme imposée par la Communauté européenne — pourrait remettre en question l'ordre établi. De plus, il n'est pas certain que l'UTP 100 Ohms puisse fonctionner au-delà du gigabit.

En conclusion, le câble UTP convient pour des réalisations de petite taille ; le câble STP (avec un écran collectif) est le minimum conseillé pour bâtir un réseau évolutif vers les hauts débits ; et le SFTP (avec un écran par paire et un blindage collectif) offre une garantie supplémentaire si vous disposez du budget nécessaire.

Catégories 5, 6 ou 7 ?

Actuellement, la catégorie 5 (100 MHz de fréquence maximale) est la plus répandue et supporte le Gigabit Ethernet sur ses quatre paires (250 Mbit/s par paire). La catégorie 6 permet de doubler le débit (250 MHz) et sa normalisation est stable. En revanche, le câble catégorie 7 n'a toujours pas de connecteur normalisé.

Câble	Caractéristiques	État de la norme
Cat 5 / Classe D	100 MHz, RJ45, 10bT, 100bT et Gigabit si testé TSB-95	TIA/EIA-568A
Cat 5E / Classe D	100 MHz, RJ45 jusqu'à 1 gigabit	TIA/EIA-568A-5 Ratification février 2000
Cat 6 / Classe E	250 MHz, RJ45 jusqu'à 2,5 Gbits au moins	RJ45TIA/EIA-568A-6 Ratification début 2001
Cat 7 / Classe F	600 MHz, prise non définie jusqu'à 10 Gbit/s	Spécifications en cours

En conclusion, le câblage catégorie 5, actuellement le plus répandu, supporte le 100bT et même le Gigabit à condition qu'il soit testé selon de nouveaux critères (norme TSB-95), tels que la paradiaphonie cumulée. Pour de nouvelles installations, on préférera donc des câbles certifiés catégorie 5E supportant d'entrée de jeu le Gigabit, voire certifiés catégorie 6 si le budget le permet.

Par ailleurs, le câble cuivre qui convient à tous les usages actuels et pour lequel les équipements actifs sont les plus répandus est le 100 Ω . Le choix du UTP, du FTP, du STP ou du SFTP dépend des perturbations électromagnétiques rencontrées dans l'immeuble (éclairage, transformateurs, moteurs, etc.), mais le SFTP est plus à même de répondre à des besoins futurs grâce à la protection maximale qu'il offre contre les perturbations.

Le choix de la fibre optique entre les locaux techniques

Généralement, le choix de la fibre optique se justifie essentiellement pour des questions de distance, au-delà de la limitation à quatre-vingt-dix mètres de la paire torsadée.

LA COMPATIBILITÉ ÉLECTROMAGNÉTIQUE (EMC)

Quand un câble est exposé à un champ électromagnétique normal, un **courant est induit** sur chacune des paires du câble en cuivre. Sa puissance varie entre **1 et 50 mv** pour un câble UTP catégorie 5, et entre **0 et 0,5 mv** pour un câble FTP. Elle dépend de la qualité du câble (mais pas de sa longueur) ainsi que de la fréquence et de la puissance du signal perturbateur. Pour les Américains, promoteurs de l'UTP, l'enjeu est énorme : trouver des parades ou changer de câble !

Il est à noter que l'inverse est vrai : le câble ne doit pas rayonner au point de générer des interférences sur les autres équipements électroniques.

Vous pouvez vous-même faire l'expérience de ce phénomène avec le tuteur de votre chaîne hi-fi : même débranché, le câble qui le relie à l'ampli est capable de capter des émissions radio de manière suffisamment puissante pour activer ce petit haut-parleur. Il est même possible d'entendre la radio en branchant un simple écouteur téléphonique au niveau du panneau de brassage !

Jusqu'à présent, ce phénomène ne perturbait pas les réseaux locaux, mais, de nos jours, les fréquences utilisées avoisinent les 100 MHz, ce qui correspond très exactement à la gamme de fréquences des radios FM et des talkies-walkies.

Aujourd'hui, les câbles catégorie 7 sont prévus pour fonctionner jusqu'à 600 MHz, et il est probable que, dans le futur, les fréquences continuent d'augmenter pour avoisiner celles du téléphone DECT (1 800 MHz) et du GSM (900 et 1 800 MHz).

Il est donc important que les câbles aient une bonne performance EMC évaluée en mesurant l'**atténuation de couplage** (AC). Les valeurs précises sont en cours de normalisation :

- AC < 40 dB : la qualité EMC est mauvaise ;
- 41 < AC < 50 : minimum requis pour les câbles UTP catégorie 5 (100 MHz) ;
- 51 < AC < 60 : minimum requis pour les câbles FTP catégorie 5 (100 MHz) ;
- 61 < AC < 70 : minimum requis pour les câbles FTP catégorie 6 (200 MHz) ;
- 71 < AC < 80 : bon câble : FTP ; câble médiocre : SFTP ;
- 81 < AC < 90 : minimum requis pour les câbles FTP et SFTP catégorie 7 (600 MHz).

Ainsi, le Gigabit Ethernet peut fonctionner sur un câble UTP catégorie 5 en respectant l'EMC uniquement si l'AC est supérieur à 50 dB.

Cependant, un système de câblage en fibre optique coûte quasiment le même prix qu'un équivalent en cuivre catégorie 7. Le choix de ce support peut donc être pris en considération pour la distribution, d'autant plus que la connectique en bureau, de types SC ou MT-RJ, est désormais de bonne qualité.

La fibre optique offre également un gage de pérennité pour le support des hauts débits.

Multimode ou monomode ?

Une fibre monomode offre de meilleures performances mais coûte plus cher que la multimode, de même que les cartes réseaux correspondantes. De plus, la multimode convient à presque tous les usages pour la mise en place de réseaux locaux au sein d'un bâtiment ou sur un campus.

Toutefois, si les distances sont réellement importantes, le choix de la monomode s'impose. Elle pourra être envisagée pour connecter d'autres sites à hauts débits dans le cadre d'une boucle optique sur un campus ou avec un opérateur.

62,5/125 ou 50/125 ?

Cette question porte sur les diamètres du cœur et de la gaine optique de la fibre, exprimés en microns. Un cœur de plus petit diamètre affaiblit moins le signal et permet donc de le véhiculer sur de plus grandes distances.

Un câble...	... supporte le Gibabit sur
Cuivre cat 5, 5E, 6, 7	90 mètres
Multimode 62,5/125	300 m à 850 nm 550 m à 1 300 nm
Multimode 50/125	550 m à 850 nm et à 1 300 nm
Monomode	3 km à 1300 nm

Cependant, la fibre optique qui convient à tous les usages actuels et pour laquelle les équipements actifs sont le plus répandus est la multimode 62,5/125. C'est également la moins chère. Tous les réseaux Ethernet, du 10 Mbit/s au gigabit, fonctionnent avec une longueur d'onde de 850 nm ou, plus rarement, de 1 300 nm.

Le câble contenant les fibres

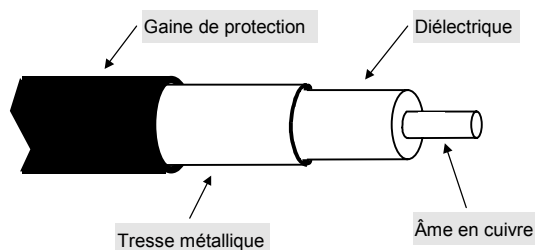
Les qualités d'un câble en fibre optique dépendent de ses caractéristiques optiques mais aussi mécaniques : gaine rigide ou souple, traitée anti-rongeurs et isolations (thermique, incendie, corrosion) qui conditionnent sa durée de vie (dix à quarante ans).

Au sein d'un bâtiment, il conviendra de choisir un câble dit d'intérieur, souple, tandis que, pour les connexions entre bâtiments, on choisira un câble dit d'extérieur, plus rigide, dont la gaine extérieure peut même être métallique.

Le coaxial et la paire torsadée pour la vidéo

Le mode de diffusion le plus répandu pour la vidéo est actuellement le câble coaxial (organisation en bus), car la plupart des équipements sont pourvus de ce type de prise. Un seul câble parcourt alors tout l'immeuble et véhicule plusieurs dizaines de canaux vidéo.

Figure 5-3.
Le câble coaxial.

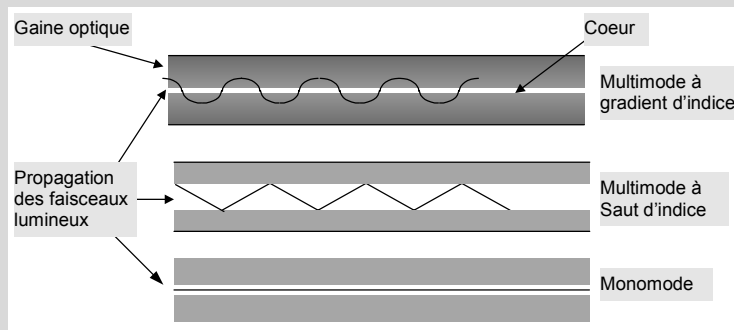


LE POINT SUR LA FIBRE OPTIQUE (EIA/TIA 492AAAA)

Ce type de câble véhicule des ondes lumineuses au sein d'une fibre caractérisée par sa **gaine optique** et son **cœur**, se différenciant par leur indice de réfraction.

La qualité d'une fibre dépend de trois paramètres :

- du mode de **propagation** de la lumière (**multimode** ou **monomode**) ;
- du **diamètres** de la gaine et du cœur (62,5/125 μ ou 50/125 μ pour la multimode) ;
- de leur **composition** (verre de silice, plastique ou composite).



La qualité du signal dépend en plus :

- de la **longueur d'onde** émise (850 et 1 300 nm pour les multimodes ; 1 310 et 1 550 nm pour les monomodes) ;
- de la source lumineuse : une diode électroluminescente LED (*Light-Emitting Diode*) ou laser ILD (*Injection Laser Diode*).

Les tableaux suivants présentent les performances comparées en fonction de ces paramètres.

	Multimode à saut d'indice	Multimode à gradient d'indice	Monomode
Source lumineuse	LED ou laser	LED ou laser	Laser
Bande passante	20 à 200 MHz/km	200 MHz à 1,5 GHz/km	3 à 50 GHz/km
Diamètre du cœur	de 50 à 125 μ	de 50 à 125 μ	de 2 à 8 μ
Diamètre de la gaine	de 125 à 440 μ	de 125 à 440 μ	de 15 à 60 μ

Type de fibre	Composition Cœur / Gaine	Affaiblissement en dB/km		
		À 850 nm	À 1 300 nm	À 1 500 nm
Multimode à saut d'indice	Verre de silice / verre de silice ou plastique ou verre composite	2	0,5	0,2
		2,5	---*	---*
		3,4	---*	---*
Multimode à gradient d'indice	Verre de silice / verre de silice ou verre composite	2	0,5	0,2
		3,5	1,5	---*
Monomode	Verre de silice / verre de silice	2	0,5	0,2

*--- Affaiblissement trop important, non utilisable

En revanche, la diffusion sur câbles de cuivre à paires torsadées (organisation en bus-étoile) tend à se généraliser, car elle permet de banaliser le système de câblage et donc de profiter de sa souplesse en termes de reconfiguration et d'évolutivité.

Enfin, la diffusion vidéo sur IP (norme H.323) fait désormais partie de l'offre des constructeurs de matériels vidéo tels que Tonna (gamme de produits Viscable++). Or, qui dit IP dit réseau local Ethernet et donc paire torsadée.

Sur quels critères choisir le type de câble ?

Dans le premier cas, la connexion aux équipements est simple : en bureau, un cordon coaxial relie la prise à une télévision ou à une carte dans un PC (de type WinPC) ; à son extrémité, le câble est connecté à la régie vidéo située dans le local nodal.

Dans le second cas de figure, la connexion est plus coûteuse, car elle requiert l'installation d'équipements intermédiaires dans les LTE. L'architecture ressemble alors à celle mise en place pour le réseau local et la téléphonie.

Il est envisageable de connecter directement la régie vidéo aux prises utilisateur en brassant les prises de distribution aux câbles de rocades jusqu'à l'endroit où est située la régie. Cette solution nécessite néanmoins beaucoup de câbles en cuivre.

En définitive, le câble coaxial est adapté à des besoins ponctuels de diffusion vidéo (moins d'une centaine de postes de travail, des salles de conférence, etc.), tandis que le câble à paires torsadées est bien mieux adapté à des gros besoins, tels que ceux nécessités dans le monde de l'audiovisuel.

Par contre, avec la généralisation de la vidéo sur IP, le câble coaxial risque bien de disparaître au profit de la paire torsadée.

Quel type de prise ?

Les standards étant bien établis, les choix sont ici plus limités :

- **RJ45** pour la paire torsadées (la prise RJ11 du téléphone peut s'insérer dans une prise RJ45 femelle, mais pas l'inverse) ;
- **SC** pour la fibre optique (attention, on trouve encore du ST) ;
- **BNC** pour le coaxial.

Pour des questions de simplicité, on met toujours le même type de prise en bureau et en local technique. Cela permet d'utiliser les mêmes cordons de brassage.

L'aménagement des locaux techniques

L'aménagement d'un local technique est plus complexe qu'il n'y paraît, car il n'y a pas de solution universelle. Ce qui doit présider à sa conception est la facilité d'utilisation, à savoir l'accès aux équipements actifs et la facilité de brassage.

Ce qui complique la tâche, c'est que les locaux dédiés à l'informatique sont généralement de petite dimension. Ils doivent néanmoins accueillir le câblage d'étage ainsi que les équipements actifs.

Les baies

La hauteur utile d'une baie est généralement de 36 ou **42 U** (un U équivalant à 4,44 cm) ; ses largeur et profondeur peuvent varier entre 800 × 800 cm, 600 × 800 cm ou 600 × 600 cm. La taille de 800 × 800 a ma préférence, car elle offre suffisamment d'espaces latéraux pour y faire passer des cordons de brassage et suffisamment de profondeur pour y loger tous types d'équipements actifs.

Dans tous les cas, elle doit être équipée de rails crénelés fixés sur les montants droit et gauche, de manière à offrir une largeur de **19 pouces** (48,26 cm). Les rails doivent être fixés à l'avant et à l'arrière, en retrait de **10 à 15 cm** par rapport aux façades. Cet espace permettra de fermer la porte lorsque tous les cordons de brassage seront installés. Détail pratique, mais qui est parfois oublié...

Les baies peuvent être dédiées au câblage ou mixtes câblage/équipements, accueillant, par exemple, un panneau de brassage dans leur partie haute et les équipements actifs dans leur partie basse.

Le cheminement des cordons de brassage

Si le local contient plusieurs baies, de nombreux cordons de brassage seront nécessaires pour raccorder les équipements aux panneaux de brassage : autocommutateurs, routeurs, concentrateurs, etc. Afin de maintenir une installation avec le minimum de cordons emmêlés, il est essentiel de simplifier la tâche des exploitants.

L'utilisation du faux plancher est déconseillée, car on y laisse toujours s'accumuler un sac de nœuds bien caché ; d'autre part, soulever les dalles est toujours une opération fastidieuse. Bien souvent, elles ne peuvent se soulever aisément, car il y a toujours un équipement posé dessus, à cheval entre deux dalles.

Il est, en revanche, préférable de faire circuler les cordons de brassage dans les flancs des baies équipées de guides câbles ainsi que dans un chemin de câble fixé en hauteur, à l'arrière de ces dernières (attention à ne pas en sous-estimer la largeur). Les cordons le moins souvent manipulés circuleront dans le chemin de câble, tandis que les cordons utilisés pour la distribution (connexion des prises utilisateurs aux équipements) chemineront dans les guides câbles.

Afin de faciliter les opérations d'exploitation (brassage, installation d'équipements, etc.), il convient également de réserver un dégagement de 80 cm au moins en face avant et en face arrière des baies.

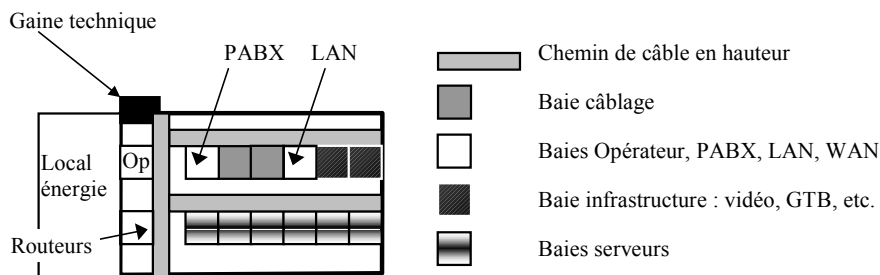
Détail non moins pratique, les luminaires seront disposés de manière à éclairer les zones de dégagement entre les baies.

L'organisation du local

Une des règles de base est de positionner les équipements actifs à proximité des panneaux de brassage, afin de limiter la longueur des cordons et également les sources de "sacs de nœuds".

Par ailleurs, tous les équipements susceptibles d'être connectés à des lignes télécom (auto-commutateurs et routeurs) doivent de préférence être situés à proximité de l'arrivée de ces lignes dans la baie de l'opérateur. Il faut dans tous les cas prévoir des câbles de déport (avec ferme CAD et/ou panneaux RJ45) entre les deux types de baies, surtout si le local opérateur est distinct du local informatique.

Figure 5-4.
*Agencement
du local nodal.*



Le local nodal peut très bien être situé dans la salle informatique ou être séparé pour des questions d'organisation, les exploitants réseaux ne devant pas avoir accès aux serveurs, et inversement pour les exploitants système.

Si les serveurs sont proches des équipements actifs, ils pourront être directement connectés à ces derniers à l'aide des cordons circulant dans les guides câbles et/ou le chemin de câble.

Si, en revanche, ils sont éloignés, ou si la salle informatique est distincte du local technique, un panneau de distribution (prises RJ45 et/ou SC) devra être installé à proximité de chaque serveur, dans une baie mixte câblage/serveur, et relié par un câblage à un panneau de distribution analogue dans le local technique.

Le cahier des charges

La rédaction de ce document synthétise les données collectées lors des phases précédentes.

Le cahier des charges a pour but d'expliquer ce qui est attendu et de guider les soumissionnaires dans leur réponse.

Dans le cadre de réalisations plus importantes, le document peut être scindé en un **CCTG** (cahier des clauses techniques générales) et un **CCTP** (cahier des clauses techniques particulières).

Le CCTG définit les engagements attendus de la part de l'entreprise :

- ses responsabilités techniques ;
- la nécessité de coordination avec d'autres corps d'état ;
- le respect d'un calendrier de réalisation ;
- le maintien du site dans un bon état de propreté si l'immeuble est déjà occupé, et en particulier l'évacuation des gravats à sa charge ;
- les garanties de qualité ;
- le respect du plan hygiène et sécurité.

Le CCTP doit décrire l'existant (l'infrastructure d'immeuble et le câblage, s'il existe), étayé par des plans (plan de masse, étage type, sous-sol). Il doit ensuite fournir tous les éléments techniques qui permettront aux soumissionnaires de répondre. Ses données sont en fait une synthèse de l'étude réalisée dans les phases précédentes. Il s'agit de décrire :

- le cheminement des câbles (dans ses principes généraux) ;
- les règles d'espacement par rapport aux sources de courants forts (câbles électriques, moteurs d'ascenseur, alimentations à coupure, tubes néon, etc.) ;
- les règles d'ingénierie que l'entreprise devra impérativement respecter concernant les câbles, les chemins de câble, la connexion des prises, le raccordement à la terre informatique, etc.

Dans une autre partie, le CCTP décrit les prestations attendues, à savoir :

- le percement des murs, si nécessaire ;
- la fourniture et la pose de tous les composants requis : chemins de câbles, goulottes, panneaux de brassage, baies, prises, boîtiers VDI dans lesquels viennent s'insérer les prises, etc.
- le raccordement à la terre informatique et, si nécessaire, la réalisation de la terre informatique à partir du puits de terre jusqu'à la distribution.

Des détails qui ont leur importance :

- la documentation des travaux réalisés tels que le cheminement exact des câbles et la position exacte des prises reportés sur les plans ;
- les fiches de tests de chaque prise ;
- l'étiquetage des prises avec des étiquettes gravées autocollantes (et non pas des Dimos ou du papier qui s'effacent ou se décollent au bout de quelques mois).

Et enfin, la nature des composants fournis et installés :

- pour le câblage cuivre : types de câbles de distribution et de rocares, types de connecteurs, etc. ;
- pour le câblage optique : types de câbles de distribution et de rocares, types de tiroirs optiques, etc. ;
- les types de baies : dimensions, avec ou sans portes, etc. ;
- les types de boîtiers VDI : nombre et types des prises (RJ45, CD, etc.) ;
- le descriptif des tests sur chaque prise et les valeurs à mesurer lors des tests réflectométriques (voir plus loin).

Afin de comparer facilement les réponses, un modèle de bordereau de prix tel que celui présenté ici pourra être joint.

Désignation des ouvrages	U	Qté	P.U. HT	P.T. HT
Lot 1 – Câblage cuivre (fourniture, pose et raccordement y compris toute sujétion)				
Chemins entre la salle informatique et le LN-A				
Percements	ens			
Chemins de câble	m			
Raccordement à la terre électrique	ens			

Le reproche que certains peuvent faire à une description aussi détaillée est qu'elle est du ressort de l'entreprise, que c'est son métier. Certes, mais le but est ici de bien définir le **niveau de prestation** que l'on attend. Si on ne le fait pas, les soumissionnaires feront des réponses minimales :

- avec des composants de faible qualité ;
- en omettant certains composants qui peuvent paraître accessoires, comme l'étiquetage des prises ou des colliers de fixation ;
- en calculant la longueur des câbles au plus juste (en les faisant passer en ligne droite sans respecter les contraintes d'écartement des sources de courant fort) ;
- sans prendre en compte la coordination avec les autres entreprises ;
- etc.

Ce type de réponse paraîtra attractif sur le plan financier, mais passera sous silence de nombreux aspects importants.

Il peut également être tentant de confier la réalisation du câblage à une société spécialisée en électricité. Cela serait une erreur, sauf si bien sûr elle dispose des compétences requises en courant faible. Car le câblage informatique n'a rien à voir avec l'électricité : ce sont deux métiers différents qui font appel à des expertises sans aucun rapport entre elles.

Enfin, la certification de l'entreprise pour le système de câblage proposé est un gage de qualité : non seulement elle offre la garantie du constructeur pendant dix à quinze ans mais, de plus, elle signifie que les techniciens ont suivi une formation spécifique de la part du constructeur sur le type de matériel proposé.

En fin de chantier, un représentant du constructeur vérifie (en plus de la recette dont nous parlerons plus loin) la qualité de l'installation, et appose son certificat de garantie. Celui-ci assure la remise en état, pendant dix à quinze ans, selon le constructeur, de n'importe quel composant défectueux (câble, connecteur, panneau de brassage). Si l'entreprise disparaît dans l'intervalle, le constructeur prend le relais ou désigne une autre société.

Le suivi du chantier et la recette

En cours de réalisation, il est nécessaire d'organiser un point hebdomadaire avec les représentants de la société de câblage : l'objectif est de contrôler l'avancement des travaux, de résoudre certains problèmes techniques, de préciser des détails comme le principe d'étiquetage des prises, etc.

Il est également indispensable de procéder à des visites régulières du site (une à deux fois par semaine selon l'état d'avancement des travaux). L'objectif est ici de contrôler la qualité des travaux en cours afin de procéder à d'éventuelles rectifications avant la fin du chantier. Mieux vaut détecter le plus en amont possible tout problème pouvant nécessiter la reprise des travaux supposés achevés.

L'ORGANISATION D'UN CHANTIER DE CÂBLAGE

Le **maître d'ouvrage** est le donneur d'ordre, celui qui paie, c'est-à-dire vous, le client.

Le **maître d'œuvre** est l'exécutant, le responsable des travaux ; il rend des comptes au maître d'ouvrage.

Le **soumissionnaire** est l'entreprise qui répond à l'appel d'offres ; le terme entreprise désigne l'**entreprise de câblage** qui a été retenue en tant que maître d'œuvre du projet **courants faibles**, par opposition aux câblages **courants forts** qui concernent l'électricité, généralement réalisée par une autre entreprise.

L'entreprise de câblage désigne un **chef de chantier** qui coordonne le travail des ouvriers sur site ; elle est parfois l'interlocuteur du maître d'ouvrage. Dans le cas de réalisations importantes, un **conducteur de travaux** est désigné en tant qu'interlocuteur.

De son côté, le maître d'ouvrage est souvent assisté d'un consultant qui, dans le cadre d'une intervention d'**assistance à maîtrise d'ouvrage**, assure le lien entre l'utilisateur exprimant des besoins généraux et le monde du câblage avec sa spécificité et son vocabulaire.

Une **réunion de chantier** est régulièrement organisée afin de coordonner et de suivre l'avancée des travaux. Elle a aussi pour but de coordonner les activités de l'entreprise de câblage avec d'autres **corps d'états** (électricien, société en charge de la climatisation, société en charge des faux plafonds, etc.). Sur le terrain, l'entreprise de câblage doit assurer cette coordination.

Le chantier terminé, il est nécessaire de procéder à sa recette. Celle-ci comprend la vérification exhaustive, qualitative et quantitative de l'ensemble des composants installés, ainsi que l'analyse des documents remis (cahier de tests, plans, etc.). Il s'agit notamment :

- de valider le cahier de tests fourni par l'entreprise;
- de réaliser des tests complémentaires sur un échantillon de prises avec le même réflectomètre fourni par l'entreprise de câblage ;
- d'effectuer le contrôle physique de l'installation avec le câbleur.

Si des anomalies sont constatées, la recette peut donner lieu à des réserves. Les réserves sont levées seulement lorsque l'entreprise de câblage a corrigé les défauts. Le procès verbal de recette peut alors être signé.

La fibre optique doit également faire l'objet d'un test réflectométrique dans les deux sens et aux deux longueurs d'ondes de référence.

La plupart des équipements réseau utilisent, en effet, la longueur d'onde de 850 nm. Mais il n'est pas exclu qu'avec le 10 Gigabit, la longueur de 1 300 nm soit la seule possible. Il faut donc tester les fibres avec les deux longueurs d'ondes et dans les deux sens !

Architecture des réseaux locaux

La conception d'une architecture réseau est élaborée en fonction du nombre de postes de travail à connecter. On peut considérer les cas suivants :

- petit réseau : moins de 200 postes dans un même bâtiment ;
- réseau moyen : de 200 à 800 postes dans un même bâtiment ;
- gros réseau : plus de 800 postes dans un même bâtiment ;
- 1^{ère} variante : plusieurs bâtiments contenant un nombre varié de postes de travail ;
- 2^{ème} variante : plusieurs sites contenant un nombre varié de postes de travail.

Bien qu'arbitraires, ces bornes correspondent à des ordres de grandeur et à des sauts technologiques. En effet, plus le nombre de postes est important, plus il faut répondre à un certain nombre de contraintes et d'exigences qui n'apparaissent qu'avec la complexité du réseau.

L'architecture réseau est bien sûre liée au système de câblage, mais celui-ci doit avoir été conçu pour en limiter les contraintes, c'est-à-dire s'adapter à toutes les situations. Le chapitre précédent a aidé à œuvrer en ce sens. Tous les systèmes de câblage sont en étoile, de même que la topologie ; les équipements actifs seront donc positionnés dans les locaux techniques. Ils serviront à connecter les postes de travail aux serveurs.

Dans le cas le plus simple, la conception d'une architecture consiste à choisir et à positionner les équipements actifs, puis à les connecter entre eux en utilisant le câblage.

Les choix de base

Le système de câblage étant maintenant prêt à l'usage, on peut alors commencer à y installer un réseau.

Quel type de réseau choisir ?

La création d'un réseau local nécessite de faire des choix, et tout d'abord celui du type : Ethernet, Token-Ring ou un autre ?

Le premier est le plus répandu et le moins cher, alors que, pour le deuxième c'est l'inverse. Le choix ira donc de préférence au premier. L'intérêt de Token-Ring est surtout sa compatibilité avec les équipements grands systèmes IBM (3090, etc.). Si vous avez des contrôleurs 3174, des terminaux 3270 ou des PC devant se connecter en émulation 3270 à un système central IBM, le choix se portera naturellement vers ce type de réseau local. Mais ce n'est pas une obligation, car la plupart des équipements IBM acceptent désormais des cartes Ethernet. Cela étant, les postes de travail peuvent toujours être connectés à un réseau Ethernet, et les systèmes IBM à des réseaux Token-Ring.

D'autres solutions sont envisageables pour construire un réseau local, mais elles sont nettement plus chères. Il s'agit, par exemple, d'ATM (*Asynchronous Transfert Mode*). Ce type de réseau est surtout destiné à d'autres usages que nous verrons au chapitre 10.

Pour toutes ces raisons, le choix d'Ethernet s'impose.

Concernant la topologie, la plus pratique est celle de l'étoile : tous les systèmes de câblage sont, on l'a vu, fondés sur ce principe.

Quel débit retenir ?

La décision suivante concerne le débit du réseau, c'est-à-dire la vitesse de transmission des trames Ethernet, encore appelée bande passante. La norme Ethernet est déclinée en plusieurs variantes : 10 Mbit/s (norme 10bT), 100 Mbit/s (norme 100bT) et 1 Gbit/s (norme 1000bT). De par son coût et son caractère innovateur, le Gigabit Ethernet est réservé aux liaisons entre les équipements de concentration et aux serveurs.

LES DIFFÉRENTS RÉSEAUX ÉTHERNET

Il existe aujourd'hui trois déclinaisons d'Ethernet normalisées par l'IEEE (*Institute of Electrical and Electronics Engineers*) : le **10bT** à 10 Mbit/s (norme 802.3), le **100bT** alias Fast Ethernet à 100 Mbit/s (norme 802.3u) et le **1000bT** alias Gigabit Ethernet (norme 802.3ab).

Le premier chiffre qualifie le débit du réseau Ethernet, la lettre "b" signifie un codage des signaux en bande de base (codage Manchester) et la lettre "T" représente "Twisted Pair", ce qui signifie que le réseau Ethernet fonctionne sur un câblage en cuivre paires torsadées.

Il existe également les mêmes déclinaisons fonctionnant sur fibre optique : **10bF**, **100bF** et **1000bX** (norme 802.3z). Parmi cette dernière, on distingue le **1000bSX** (S pour *short wavelength*) opérant à 850 nm sur fibre optique multimode et le **1000bLX** (L pour *long wavelength*) opérant à 1 300 nm sur les fibres multimode et monomode.

Ces variantes utilisent toutes le même principe d'accès au support de transmission (détection de collision), les mêmes trames Ethernet et le même adressage MAC (*Medium Access Control*).

Il existe différents types d'équipements : les **concentrateurs** (*hubs*), qui partagent un segment Ethernet entre plusieurs ports, et les **commutateurs** (*switch*), qui créent un segment Ethernet par port.

Le concentrateur peut être segmenté en plusieurs réseaux indépendants (cela dépend du modèle), tandis que le commutateur est capable d'**interconnecter** les réseaux qui sont physiquement indépendants.

Une carte 10bT d'un PC ne peut être connectée qu'à un hub ou à un switch 10bT. Un réseau Ethernet ne peut pas mélanger les débits.

Le choix du débit se fera donc en fonction des coûts, plutôt 10/100 Mbit/s pour les PC et 100/1000 Mbit/s pour les serveurs.

Débit	Utilisation
10 Mbit/s	Postes de travail bureautique
100 Mbit/s	Postes de travail multimédias et serveurs
1 Gbit/s	Pour connecter les équipements réseaux entre eux ainsi que les gros serveurs

Côté poste de travail, la plupart des cartes réseau fonctionnent à 10 et 100 Mbit/s, et sont au même prix que les cartes 10 Mbit/s. En outre, la plupart des commutateurs offrent des ports à détection automatique de vitesse (port *autosense*).

Quel format d'équipement ?

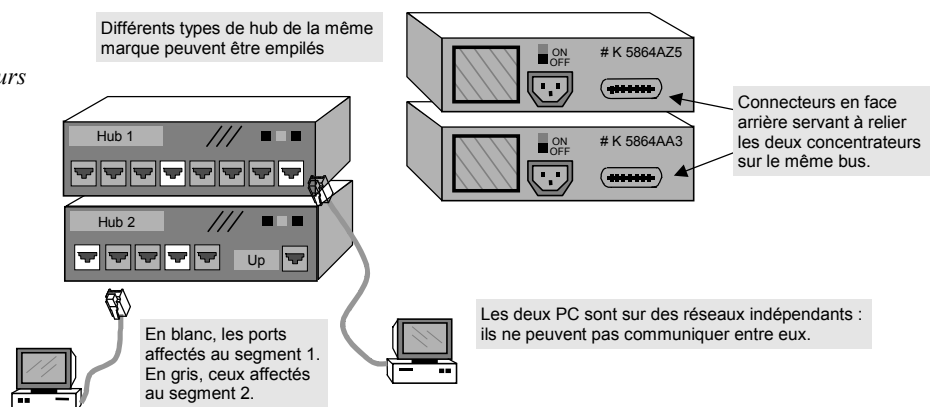
Le marché offre le choix d'équipements seuls (*stand alone*), empilables (*stackable*), ou en châssis : des cartes sont insérées dans des emplacements prévus à cet effet (*slots*). Pour compliquer le choix, il existe également des *stackables* avec des *slots* d'extension qui permettent d'ajouter un ou deux ports, dans le but de chaîner l'équipement à un autre.

Les modèles *stand alone* visent le marché d'entrée de gamme ; ils sont parfaits pour créer un premier réseau local (voir chapitre 6).

Les modèles empilables sont envisageables dès qu'il y a des possibilités d'extension. Par exemple, votre société dispose de cinquante postes de travail, et vous commencez par en connecter dix dans un premier temps. Vous pouvez alors acheter un hub 12 ports, puis un autre 12 ports plus tard.

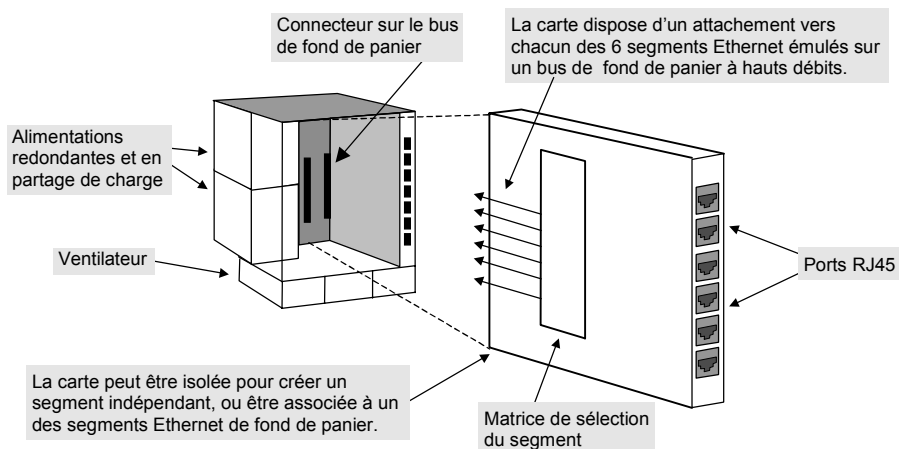
Un concentrateur, ou un commutateur, comprend un nombre limité de ports (généralement 8, 16, 24 ou 32). Or, la plupart du temps, un local technique concentre beaucoup plus de postes de travail (jusqu'à plusieurs centaines). La solution est alors de chaîner les *stackables* entre eux *via* un bus spécial, dédié à cet effet ; il s'agit d'un câble externe reliant les équipements entre eux pour n'en faire qu'une unité logique. Attention, le bus étant propriétaire, seuls les équipements d'un même constructeur pourront être chaînés entre eux, généralement de 5 à 8 au maximum.

Figure 6-1.
Les possibilités
des concentrateurs
empilables.



Une autre solution consiste à installer des châssis, certes plus chers, mais qui offrent de plus grandes capacités d'accueil. Ces équipements permettent de créer plusieurs segments indépendants à des débits différents.

Figure 6-2.
*Principes
et fonctionnalités
d'un châssis.*



Il est possible d'insérer différents types de cartes dans un châssis : concentrateur Ethernet, commutateur Ethernet, carte Token-Ring, FDDI, ATM, etc. Il est également possible de combiner les débits (10, 100 et 1 000 Mbit/s) sur des segments séparés. Mais, attention, il n'est en aucun cas possible de mélanger des débits sur un même segment Ethernet. De même, les réseaux Token-Ring, FDDI et ATM créés seront indépendants.

Le tableau suivant présente quelques éléments de comparaison.

Critère	Empilable	Châssis
Évolutivité	Ajout d'éléments empilables limité à cinq environ	Ajout de cartes limité par le nombre de slots
Segmentation	Limitée à un ou deux segments	Plusieurs segments par port ou par groupe de ports
Capacité de traitement	Bus externe limité à quelques centaines de Mbit/s	Bus de fond de panier de 100 Mbit/s à plusieurs Gbit/s
Alimentation	Une par élément ou, rarement, une pour tous	Une à trois pour l'ensemble du châssis
Redondance d'alimentation	Pas tout le temps	Oui
Création de plusieurs segments indépendants	Oui	Oui
Utilisation	Moins de cent postes en Ethernet, en Token-Ring ou en ATM	Plus de cent postes ; combinaison Ethernet, Token-Ring et ATM possible

En définitive, l'empilable sera choisi pour des faibles densités (moins de cent postes par local technique) ; les châssis seront privilégiés dans les autres cas. D'une manière générale, plus le réseau concentre de postes de travail, plus la fiabilité des équipements doit être importante.

Les châssis seront donc choisis là où le besoin en bande passante est élevé et où la fiabilité est primordiale, c'est-à-dire à des points de concentration stratégiques du réseau.

Concentrateur ou commutateur ?

L'autre décision à prendre consiste à choisir entre les concentrateurs et les commutateurs. Les premiers se contentent de générer le signal, alors que les seconds permettent de créer un segment par port. Ces derniers sont bien sûr plus chers.

Pour une utilisation bureautique du réseau (traitement de texte, comptabilité, base de données, connexion à un serveur, etc.), les concentrateurs suffisent pour connecter les postes de travail car il y a peu de trafic entre eux. Pour améliorer les performances, on peut jouer sur la vitesse (10 ou 100 Mbit/s).

L'utilisation des commutateurs s'envisage dans plusieurs cas de figures :

- lorsqu'on emploie des applications multimédias (voix et vidéo) générant des débits importants et nécessitant des temps de réponse courts ;
- d'une manière générale, lorsque le flux réseau est important et que les temps de réponse sont mauvais ;
- pour interconnecter plusieurs segments Ethernet.

Si vous constatez un nombre élevé de collisions lié à une charge réseau importante, vous pouvez, dans un premier temps, segmenter le réseau (c'est-à-dire le couper en deux). Dans ce cas, les PC situés sur un segment ne pourront plus communiquer avec ceux situés sur l'autre.

Au lieu d'acheter un second concentrateur, l'achat d'un commutateur résoudra le problème : les postes seront répartis sur les deux équipements, et les segments ajoutés seront interconnectés.

QU'EST-CE QU'UN COMMUTATEUR ?

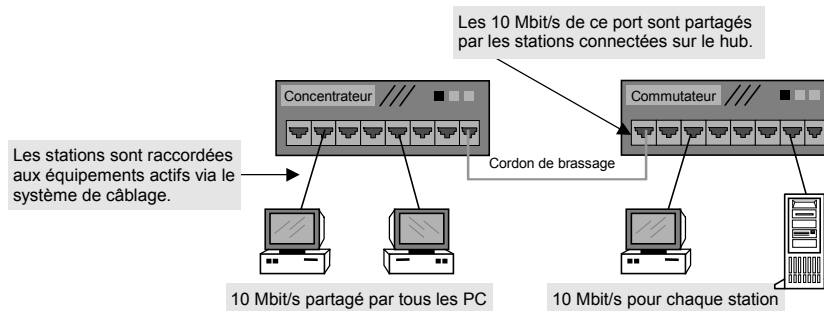
Un commutateur (*switch*) est un équipement qui offre une bande passante dédiée pour chaque port (10, 100 ou 1 000 Mbit/s par port) alors que le concentrateur partage la bande passante entre tous ses ports. Cela revient à créer un segment Ethernet par port.

On distingue les *switches cut-through* (*on the fly*, à la volée) et les *switches store and forward* (les plus courants aujourd'hui). Les premiers se contentent de régénérer la trame (même les trames erronées et les collisions), tandis que les seconds la stockent en mémoire avant de la régénérer. La méthode *adaptive cut-through* combine les deux principes : dès qu'une trame est en erreur, le commutateur bascule en mode *store and forward* pendant un certain temps. La méthode *fragment free*, la plus performante, lit les 64 premiers octets avant de décider du mode de transmission.

Chaque port du commutateur correspond à un segment Ethernet, c'est-à-dire à un domaine de collision (voir encadré "Le point sur Ethernet").

Figure 6-3.

Concentrateur
et commutateur.



Les stations consommant le plus de bande passante seront de préférence connectées au commutateur. C'est le cas des serveurs qui concentrent toutes les connexions des utilisateurs.

Les commutateurs apportent, par ailleurs, de nouvelles fonctionnalités et ils peuvent être préférés aux concentrateurs rien que pour cela, indépendamment de tout problème de charge réseau.

Étant donné qu'un commutateur crée un segment par port, il est possible de combiner les débits (10, 100 et 1 000 Mbit/s) au sein de la même boîte. Il est à noter que seul le mode *store and forward* le permet tandis que le mode *cut-through*, quant à lui, ne bénéficie pas de cette possibilité car les trames sont commutées à la volée (elles ne sont pas stockées en mémoire).

On trouve sur le marché les formules suivantes :

- ports à détection automatique de vitesse (*port autosense*, 10/100 Mbit/s) ;
- port uplink à 100 Mbit/s ou 1 Gbit/s.

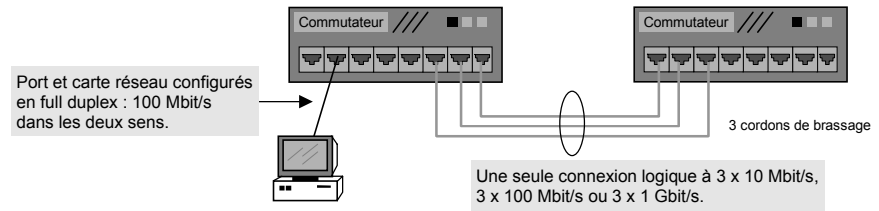


Les cartes réseau autosense 10/100 posent des problèmes avec les commutateurs qui sont également autosense, car il n'y a pas de négociation de débit ; chacun essaie de se caler sur la vitesse de l'autre. Il est donc recommandé de désactiver cette fonction au niveau de la carte et de fixer la vitesse manuellement (configuration à l'aide de Windows).

Les commutateurs permettent également d'augmenter les débits de plusieurs manières :

- avec le mode full duplex entre un PC et un port du commutateur ou entre deux commutateurs ;
- en agrégeant plusieurs ports full duplex du commutateur (technique du *port trunking*) pour le relier à un autre commutateur.

Figure 6-4.
*Augmentation
des débits
grâce
aux commutateurs.*



LE POINT SUR L'ETHERNET FULL DUPLEX (IEEE 802.3x)

Le mode full duplex **conserve le débit nominal** de 10, 100 ou 1 000 Mbit/s, mais **sépare les canaux émission et réception**, et donc, **élimine le besoin de la détection de collision CSMA/CD** employée par l'Ethernet classique (en half duplex).

Le mode full duplex n'est possible :

- que sur des liaisons point à point entre un PC et un commutateur ou entre deux commutateurs ;
- qu'avec des câbles qui séparent physiquement les canaux émission et réception, c'est-à-dire le cuivre à paires torsadées et la fibre optique, mais pas les câbles coaxiaux ;
- qu'avec les commutateurs qui sont seuls capables de stocker les données à envoyer lorsque le canal émission est occupé.

L'absence de détection de collision CSMA/CD a une conséquence importante : **la limitation de longueur des câbles imposée par le délai de propagation des collisions disparaît**. Seule l'atténuation du signal, et donc la performance du câble, limite désormais la distance entre deux équipements : en pratique, on atteint des valeurs comprises entre 150 et 200 m au lieu des 100 m habituels, et plusieurs centaines de kilomètres en fibre optique.

Le mode full duplex ne permet pas d'augmenter la vitesse de transmission, mais consiste à séparer les canaux émission et réception. Le débit global est ainsi augmenté, et peut, théoriquement, être multiplié par deux. Dans la pratique, ce mode est donc intéressant pour les serveurs et les liaisons intercommutateurs qui peuvent avoir à traiter un flux simultané dans les deux sens.

La carte réseau doit également supporter le mode full duplex.

À l'inverse, l'agrégation consiste à créer une seule liaison logique constituée de plusieurs liaisons physiques. Le débit obtenu est alors égal à la somme des débits des ports agrégés. Cette technique, normalisée 802.1q, est utilisée pour interconnecter deux commutateurs généralement de même marque. Certaines cartes réseau supportent ce mode de fonctionnement.

Fonctionnalité	Concentrateur	Commutateur
Segment	Un seul pour le concentrateur	Un par port
Segmentation	(Optionnel) Réseaux indépendants (c'est-à-dire isolés, ne pouvant pas communiquer entre eux)	Segments interconnectés par la matrice de commutation
Mélange des débits	Non	Oui pour le mode <i>store and forward</i>
Port Uplink	Pour chaîner les hubs ; 4 maximum en cascade	Pour chaîner les commutateurs entre eux
Administrable	SNMP en option	SNMP + RMON en option
Autres fonctionnalités	Partitionnement des ports sur erreur	Détection automatique 10/100 Mbit/s Full duplex Agrégation de ports VLAN

Les VLAN (*Virtual LAN*) seront étudiés au chapitre 11. Pour l’instant, nous n’en avons pas besoin.

L’architecture

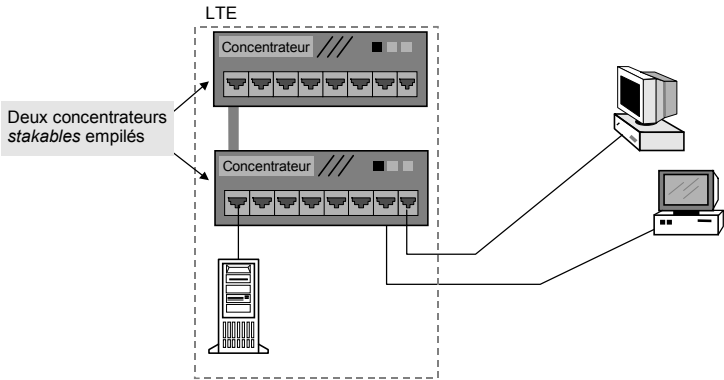
Mise en place d’un réseau local d’étage

Ayant en tête toutes les possibilités des équipements à notre disposition, la conception d’une architecture réseau simple consiste à assembler les concentrateurs et les commutateurs en exploitant au mieux les capacités du câblage.

Partons d’un cas simple : une cinquantaine de PC situés au même étage d’un immeuble quelconque. Les utilisateurs ont juste besoin d’échanger des données entre eux et de partager des applications (un serveur web, une base de données, des traitements de texte, etc.).

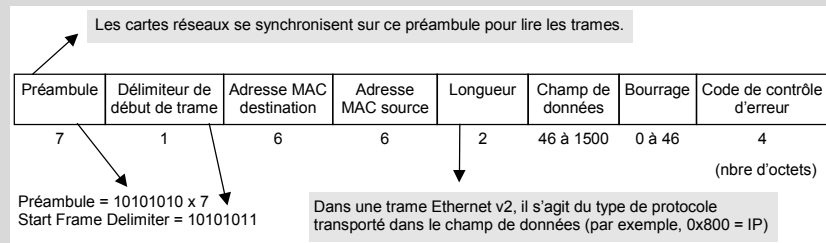
Le schéma suivant décrit l’architecture de base qui en résulte.

Figure 6-5.
Réseau local
sur un étage.



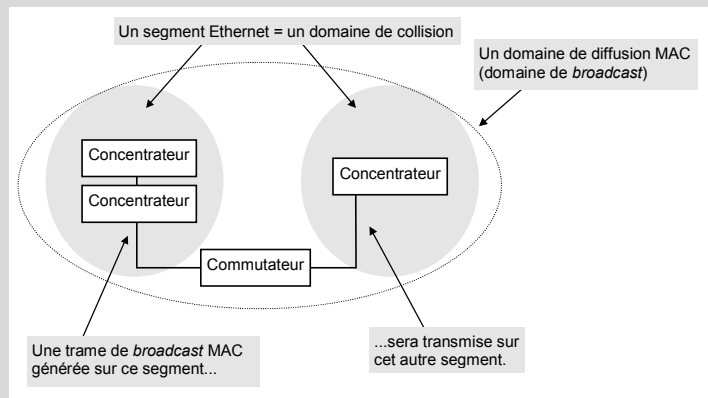
LE POINT SUR ÉTHERNET (IEEE 802.3)

La norme Ethernet décrit un protocole de niveau 1 (couche physique) — c'est-à-dire la manière de générer les signaux sur le câble — ainsi que le protocole de niveau 2 (couche liaison appelée également couche **MAC**, *Medium Access Control*). La couche MAC traite une série de bits structurés sous forme de trame, la **trame MAC**.



Chaque carte réseau est identifiée par une adresse MAC unique sur 6 octets (par exemple, 01:00:0A:FB:5D:52). Toute trame MAC émise comporte les adresses de l'émetteur et du destinataire. La carte réseau ne prend en compte que les trames MAC qui lui sont destinées. Exception à cette règle : une carte peut émettre une trame spéciale, appelée trame de diffusion (ou trame de **broadcast**), qui sera lue par toutes les autres cartes (adresse destination dont tous les bits sont à 1 — FF:FF:FF:FF:FF:FF).

Un **domaine de collision** comprend un ou plusieurs concentrateurs Ethernet ; une trame émise sur un port est régénérée sur tous les autres ports des concentrateurs chaînés (*via* un cordon de brassage) ou empilés (*via* un bus externe par un câble spécifique). Le commutateur ne laisse pas passer les collisions, **mais une trame de broadcast émise sur un port sera régénérée sur tous les autres ports**.



Quand un PC est connecté à un port du commutateur, c'est comme s'il était tout seul sur un segment Ethernet : aucune collision n'est donc possible, et il dispose de toute la bande passante. Inversement, tous les PC connectés à un concentrateur partagent la même bande passante (10, 100 ou 1000 Mbit/s) et peuvent émettre des trames en même temps, d'où une probabilité de plus en plus importante de collision qui croît avec le nombre de PC.

C'est on ne peut plus simple :

- En vertu de ce qui a été dit aux sections précédentes, deux concentrateurs empilables 10bT ont été installés et chaînés entre eux sur un bus externe.
- Chaque poste est raccordé à un port d'un concentrateur *via* un système de câblage tel que décrit au chapitre précédent : un câblage en cuivre à paires torsadées (avec, bien sûr, des prises RJ45) centré en étoile autour d'un local technique.
- Les serveurs sont ici situés dans le local technique, et chacun d'eux est raccordé directement à un port d'un concentrateur *via* un cordon de brassage RJ45/RJ45.

Il n'y a pas de question à se poser.

QUELS CORDONS DE BRASSAGE ?

Pour connecter un PC à un hub, un cordon de brassage **droit** (de même nature que le câblage en cuivre à paires torsadées) doit être utilisé. Pour connecter deux hubs entre eux, un cordon **croisé** (paires émission et réception) doit être utilisé, sauf si le port uplink est utilisé.

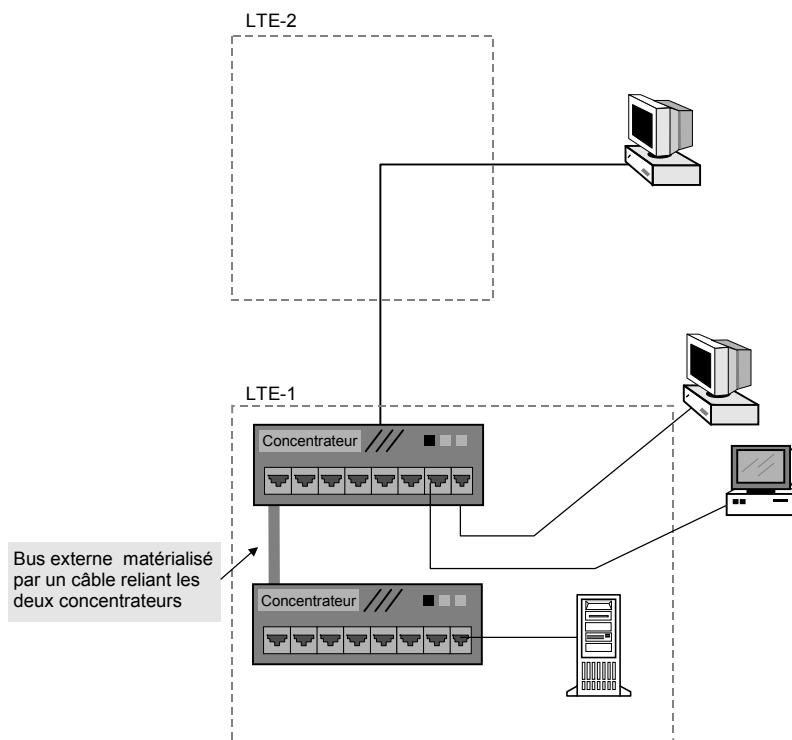
Extension du réseau d'étage

Maintenant, des utilisateurs, situés à l'étage au-dessus, ont les mêmes besoins.

Pas de problème : il y a moins de dix PC, ce qui nous permet d'utiliser les rocade en cuivre existantes, et de les connecter directement aux concentrateurs qui sont déjà installés.

Figure 6-6.

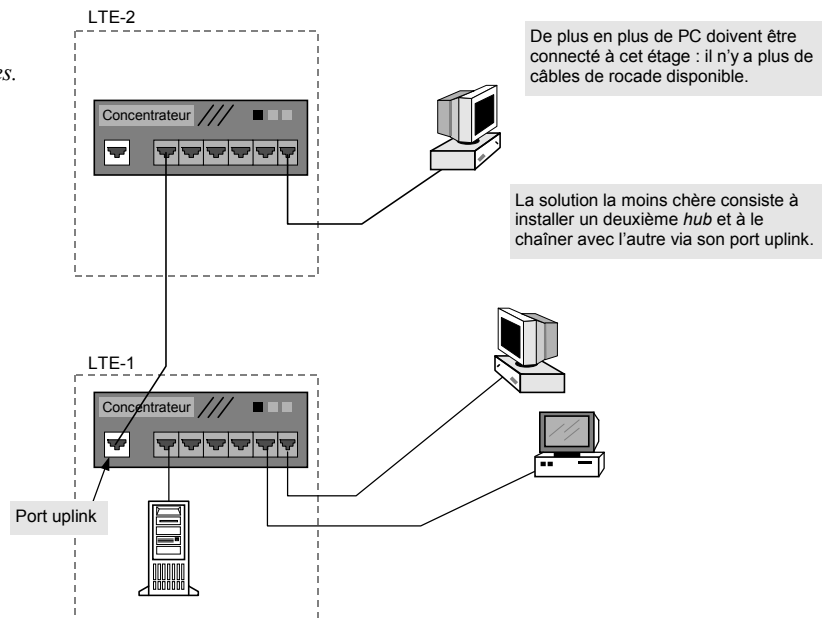
Extension d'un réseau local sur plusieurs étages.



S'il y a davantage de PC, ou, d'une manière générale, s'il n'y a plus de câbles de rcade en nombre suffisant (en général, entre 6 et 12), il suffit de créer un autre réseau sur le même modèle. Le problème est maintenant de connecter les réseaux construits sur les deux étages pour que tout le monde puisse accéder aux mêmes données et aux mêmes applications.

La solution la plus simple consiste à connecter les deux concentrateurs en cascade, *via* un câble de rcade en cuivre.

Figure 6-7.
*Un réseau local
étendu sur plusieurs étages.*



Conception d'un réseau d'immeuble

Maintenant, la situation se corse un peu : le réseau est un succès, il y a de plus en plus de demandes de connexions, la société utilise de plus en plus l'informatique. Il faut maintenant créer des réseaux à chaque étage. La solution précédente, consistant à chaîner les concentrateurs n'est plus applicable, car on est limité par le nombre de cascades possible. De plus, au-delà de cent postes connectés, le réseau Ethernet deviendrait saturé (du fait du nombre de collisions qui augmente avec le nombre d'utilisateurs) et, en définitive, les temps de réponse seraient trop grands.

L'installation de boîtes doit maintenant faire place à une plus grande réflexion, c'est-à-dire à un travail d'architecture. Imaginons donc que nous ayons trois cents utilisateurs répartis sur une demi-douzaine d'étages, soit en moyenne cinquante postes par étage, plus les imprimantes et les serveurs. On se retrouve avec soixante à soixante-dix connexions par étage.

Comme précédemment, on peut installer une pile de concentrateurs à chaque étage pour créer un réseau local d'étage. Là encore, il faut se poser de nouveau les mêmes questions :

- Quel débit : 10bT, 100bT ou Gigabit ?
- Quelle technologie : concentrateurs ou commutateurs ?
- Un seul segment Ethernet ou plusieurs ?

Les réponses à ces questions dépendent avant tout du trafic prévisionnel, des perspectives d'évolution et des performances mises en balance par rapport au coût. Si l'immeuble comprend trente autres étages, on peut supposer qu'il faudra tôt ou tard étendre le réseau. Si, en revanche, l'immeuble n'en comprend que six, on sait que la configuration sera figée pour un bon moment.

L'architecture doit donc être conçue pour couvrir les besoins futurs et non seulement ceux du moment. Elle doit donc être évolutive, c'est-à-dire être bâtie sur des équipements que l'on pourra récupérer (recycler pour d'autres usages).

Mise en place d'un réseau fédérateur

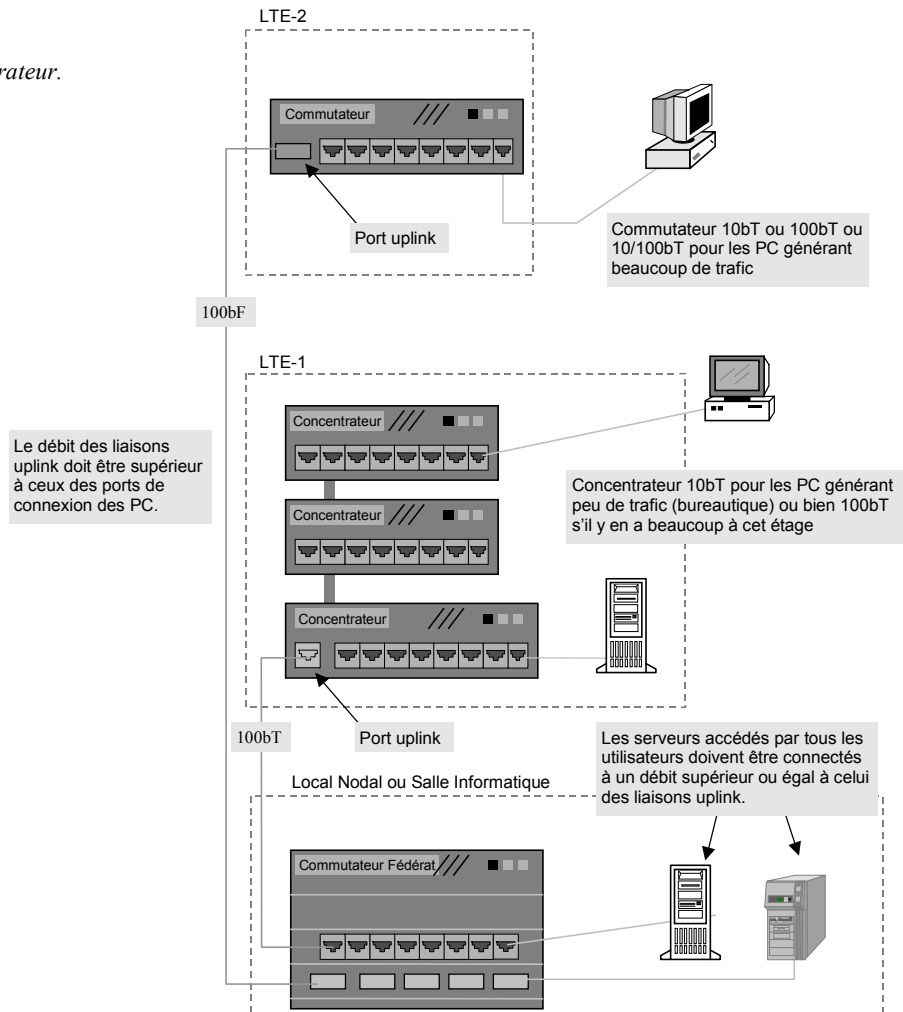
Si les besoins en trafic sont importants (applications multimédias, applications, voix, données et visioconférences), on peut envisager des commutateurs 10bT à tous les étages. Dans la plupart des cas, un débit de 10 Mbit/s suffira, mais la différence de coût étant minime, on peut envisager le 100bT.

La question est maintenant de savoir comment connecter les réseaux locaux entre eux. La solution repose sur la création d'un réseau fédérateur (*backbone*).

On peut imaginer un câble Ethernet (topologie en bus), ou FDDI (anneau), qui parcourt tous les étages et auquel on connecte les concentrateurs. C'est une solution peu évolutive, car le débit est limité à la technologie utilisée (100 Mbit/s pour FDDI). En outre, le Gigabit Ethernet et l'ATM ne sont pas prévus pour une topologie en bus ou en anneau. C'est également une solution peu sûre : le câble étant un élément passif, il n'y a aucun moyen de le superviser à distance.

L'architecture couramment utilisée est de type *collapse backbone* (littéralement, réseau fédérateur effondré). Le principe consiste à concentrer le backbone en un seul point : au lieu d'avoir un réseau qui parcourt tous les étages, le backbone est réalisé dans un commutateur unique. Cela revient à créer une architecture en étoile à deux niveaux, un premier concentrant les PC à chaque étage et un second concentrant les équipements d'étage en un point central, en général la salle informatique ou un local nodal dédié aux équipements réseau.

Figure 6-7.
*Conception
d'un réseau fédérateur.*



Pour un réseau de taille moyenne (de 200 à 800 utilisateurs), l'équipement central doit être de grande capacité en termes d'accueil et de performances. Le choix se portera donc sur un châssis qui offre une matrice de commutation à haut débit. Le réseau fédérateur n'est alors pas limité à 10 Mbit/s, mais à 100 Mbit/s par étage et à la capacité de la matrice de commutation du commutateur central, généralement plusieurs gigabits.

Quel débit et quelle technologie ?

Le choix du débit du réseau fédérateur dépend de celui utilisé par les PC.

Si les PC sont connectés à un...	Le débit des liens uplink vers le commutateur central doit être au moins égal à...
Concentrateur à 10 Mbit/s	10 Mbit/s
Concentrateur à 100 Mbit/s	100 Mbit/s
Commutateur 10 Mbit/s	100 Mbit/s
Commutateur 100 Mbit/s	1 Gbit/s

COMMENT FONCTIONNE UN COMMUTEUR ?

Un commutateur permet d'**interconnecter** plusieurs segments Ethernet. Sur chacun de ses ports, on peut raccorder un concentrateur (plusieurs PC partagent alors la bande passante sur ce port) ou un seul PC (technique de la **microsegmentation**).

Afin de limiter le trafic réseau inutile, les trames Ethernet échangées entre les PC d'un même segment ne sont pas transmises sur les autres segments gérés par le commutateur (voir encart "Le point sur Ethernet").

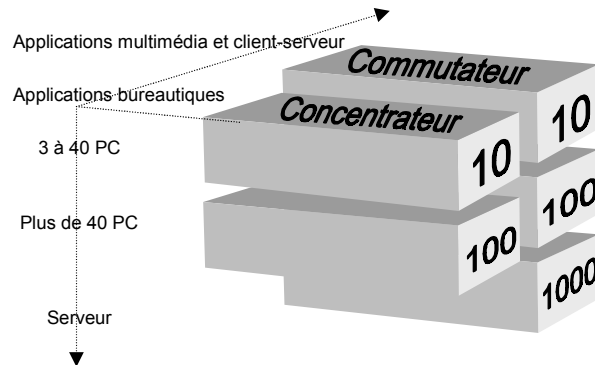
Pour savoir quelles sont les trames qui doivent sortir du segment et celles qui ne le doivent pas, le commutateur regarde toutes les adresses MAC sources des trames qui entrent sur chacun de ses ports, et les enregistre dans ses **tables d'adresses** (il existe une table par port). C'est le mécanisme d'**apprentissage**. Si un PC est déplacé d'un port à l'autre ou d'un commutateur à l'autre, son adresse MAC peut se retrouver dans deux tables et créer ainsi des conflits. Pour éviter cela, les adresses sont effacées de la table au bout de 15 à 30 secondes.

Les commutateurs d'entrée de gamme ne disposent que de très peu de mémoire et ne peuvent apprendre qu'une, deux ou quatre adresses MAC par port. Ils sont plutôt dédiés à la microsegmentation. Les commutateurs fédérateurs doivent en revanche disposer de beaucoup de mémoire et être capables d'enregistrer plusieurs milliers d'adresses MAC, car ils fédèrent tous les flux interréseaux locaux (c'est-à-dire intersegment Ethernet).

La **mémoire tampon** doit également être suffisamment importante pour permettre l'adaptation des débits (entre 10 et 100 Mbit/s, et surtout entre 10/100 et 1 Gbit/s).

L'interconnexion des segments est réalisée par une **matrice de commutation** à haut débit capable de supporter la somme des débits des ports (8×10 Mbit/s, par exemple). Là encore, les commutateurs fédérateurs doivent comprendre une matrice de commutation très puissante (généralement des ASIC et des processeurs RISC).

Le choix du débit des PC dépend, quant à lui, du volume de trafic généré et du type d'application (du flux bureautique au flux multimédia). Mais le passage du concentrateur au commutateur évite ou retarde l'augmentation du débit, ce qui permet de conserver les cartes réseau existantes dans les PC.



Indépendamment de la charge réseau et du nombre de PC, certaines combinaisons sont plus appropriées que d'autres.

Type de trafic	Réseau d'étage	Fédérateur
Applications bureautique (Word, Excel, bases de données)	Concentrateurs 10bT	Commutateurs 10/100bT
Client-serveur et un peu de multimédia (voix et vidéo)	Commutateurs 10/100bT	Commutateurs 100bT
Applications multimédias intensives	Commutateurs 100bT	Commutateurs Gigabit

Le coût est un autre critère de décision, sans doute le plus important. Lors du choix d'une technologie (concentration ou commutation) et du débit, il faut tenir compte du nombre de cartes réseau pour les PC ainsi que du nombre d'équipements.

Équipement actif 8 ports	Coût en francs HT
Carte 10bT	de 250 à 800
Cartes 100bT et 10/100bT	de 500 à 1 200
Concentrateur 10bT	de 500 à 5 000
Concentrateur 100bT (ou 10/100bT)	de 2 500 à 7 000
Commutateur 10bT	de 5 000 à 15 000
Commutateur 100bT (ou 10/100bT)	de 15 000 à 20 000
Concentrateur 1000bT	N'existe pas
Commutateur 1000bT	de 50 000 à 90 000

Les écarts de prix sont dus à des différences dans les fonctionnalités proposées (concentrateur administrable ou non, empilable ou non, avec ou sans slot d'extension, etc.). Un exemple en a été donné au chapitre 4 concernant les concentrateurs 10bT. Le choix de la fibre optique fait également monter les prix.

Suivre l'évolution des besoins

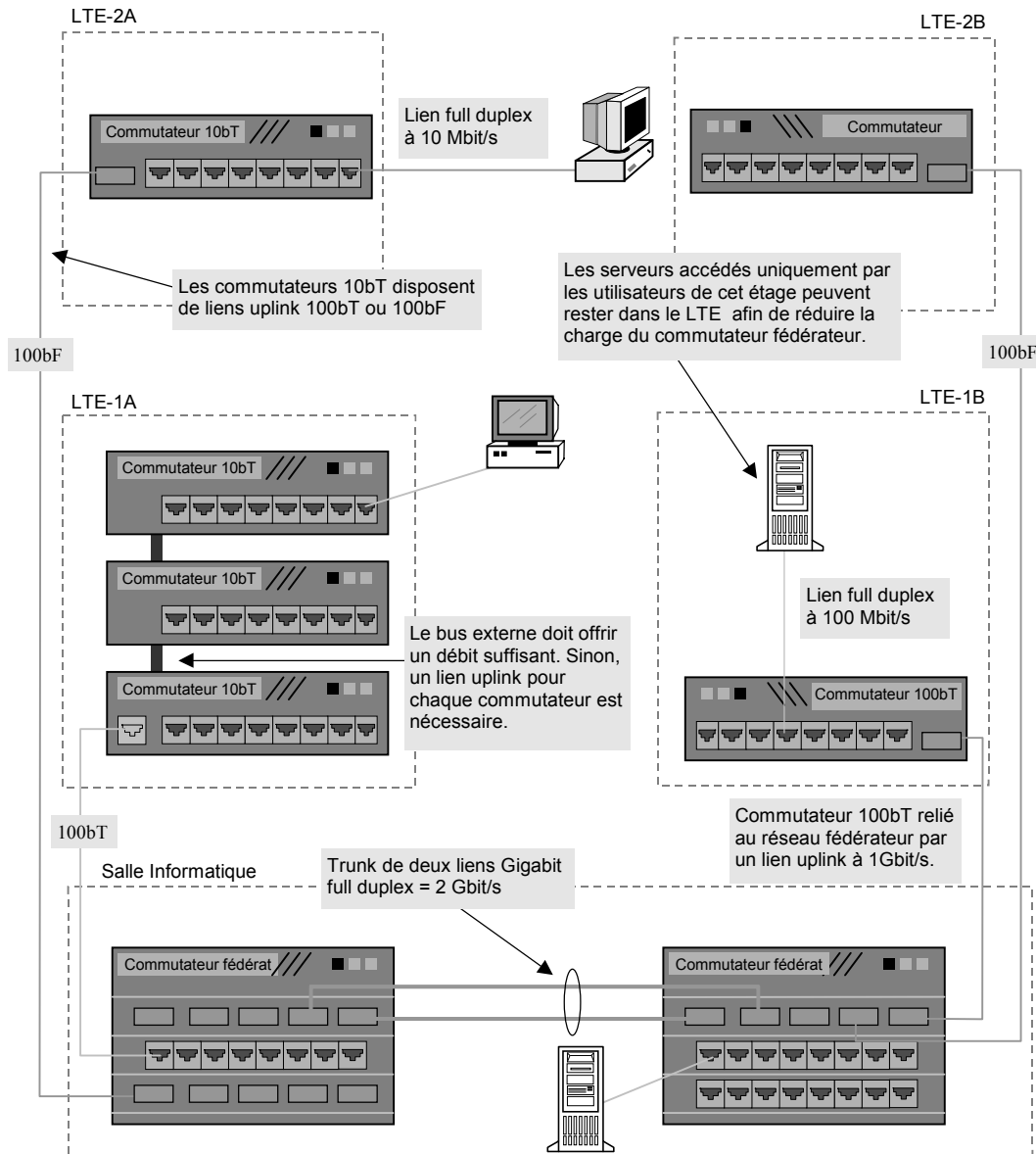
La direction a maintenant décidé de regrouper plusieurs sites sur le nôtre et de louer l'autre aile de l'immeuble. Le nombre d'utilisateurs va ainsi passer de 300 à 500. Ils seront répartis sur 10 étages, soit une moyenne de 70 connexions par étage en comptant les imprimantes, les doubles connexions, etc.

Par ailleurs, les utilisateurs de bases de données client-serveur se plaignent des mauvais temps de réponse. Il s'agit de gros consommateurs de bande passante (gestion électronique de documents, applications décisionnelles de type datawarehouse, etc.) qui perturbe les autres trafics. Les utilisateurs qui se connectent sur des serveurs Unix se plaignent également de ralentissements brusques lors de l'affichage des écrans et du déplacement du curseur à l'écran.

Pour faire face à cette montée en charge, nous avons plusieurs solutions, que nous pouvons combiner.

En premier lieu, les commutateurs peuvent être généralisés. Cela permet de segmenter le réseau afin qu'un flux important ne viennent en perturber un autre. Ensuite, les débits peuvent être augmentés à 100 Mbit/s pour les plus gros consommateurs. Afin de limiter le coût du réseau, les autres utilisateurs peuvent rester connectés à 10 Mbit/s.

Figure 6-8.
*Utilisation des capacités des commutateurs
pour monter en charge.*



Au niveau du réseau fédérateur, la capacité d'accueil doit être augmentée. La sécurité de fonctionnement doit également être améliorée, car la panne de l'équipement central entraînerait l'arrêt de tout le réseau. Le ou les commutateurs empilables doivent donc faire place à un châssis plus performant (en termes de bande passante de la matrice de commutation) et qui offre des alimentations redondantes.

L'ajout d'un second commutateur, identique au premier et donc interchangeable, renforce encore la fiabilité du réseau fédérateur.

Afin de ne pas créer un engorgement entre les deux commutateurs, un lien à très haut débit doit être créé. Cela est réalisé en agrégeant plusieurs ports Gigabit (technique du *port trunking*).

Les serveurs très sollicités et certains postes de travail consommateurs de bande passante peuvent voir leurs performances d'accès réseau accrues par des connexions full duplex.

Assurer la continuité de service

Près de 500 utilisateurs sont connectés. Le réseau devient un élément stratégique du système d'information. En effet, sans lui, les utilisateurs ne peuvent plus travailler : plus de connexion au système central, plus d'accès aux bases de données, plus de messagerie, etc. Bref, il est nécessaire de construire une architecture redondante pour parer aux éventuelles pannes. Celles-ci ne sont pas un mythe. En voici quelques-unes qui arrivent fréquemment :

- panne d'alimentation ;
- coupure d'un câble entre les commutateurs (erreur de manipulation lors du brassage dans les locaux techniques) ;
- débranchement d'un câble, d'un convertisseur ou de tout autre petit boîtier : eh oui, un *transceiver* ou un connecteur mal enfoncés dans la prise tendent non pas à se remettre en place d'eux-mêmes (dommage), mais plutôt à tomber sous l'effet des vibrations et de la pesanteur...

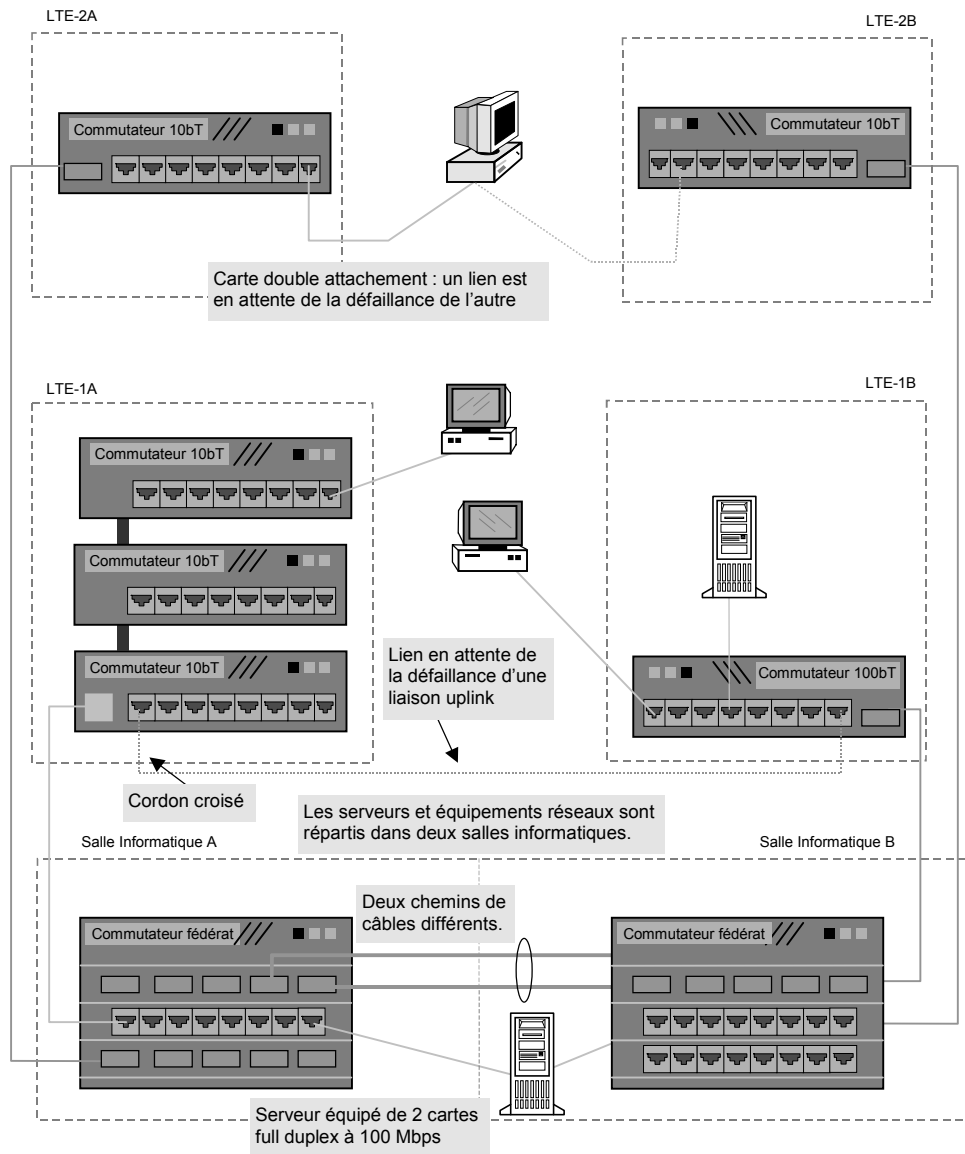
L'incendie ou l'inondation sont plus rares, mais, quand ils se produisent, il n'est pas question de reconstruire un système de câblage et un réseau avant quelques heures, voire plusieurs jours.

La mise en place d'un réseau redondant repose d'une part sur un système de câblage qui offre l'alternative de plusieurs cheminements, d'autre part sur la mise en place d'équipements de secours.

La redondance est généralement limitée au réseau fédérateur, mais, dans des cas critiques (salles de marché, processus industriels en flux tendus, etc.), il peut être nécessaire de l'étendre jusqu'au poste de travail.

On peut alors prévoir d'équiper ces derniers en carte deux ports : quand le premier perd le contact avec le port du concentrateur, il bascule sur le second, connecté à un autre concentrateur situé dans un autre local technique. Côté serveur, on peut envisager la même solution ou encore prévoir deux cartes ou plus, ce qui offre l'intérêt du partage de charge en fonctionnement normal.

Figure 6-9.
*Cas
d'un réseau
redondant.*



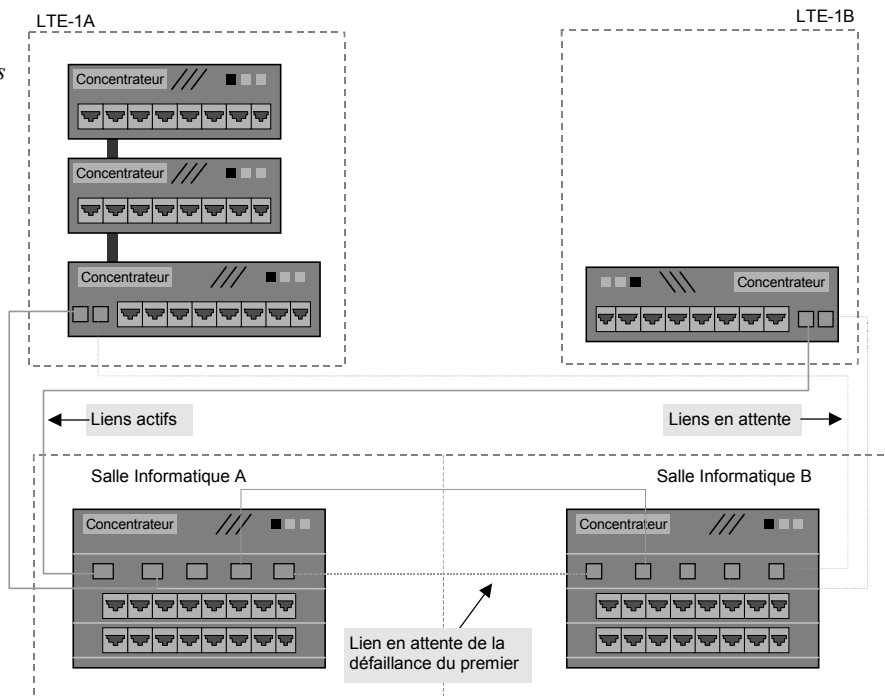
Les équipements fédérateurs sont ici des commutateurs reliés entre eux pour former un seul réseau fédérateur. Cette liaison doit être de débit élevé afin d'absorber tous les flux issus des étages et allant vers un serveur connecté au commutateur fédérateur opposé. On peut envisager ici du Gigabit Ethernet. Les liaisons entre les commutateurs fédérateurs et les équipements d'étage peuvent être réalisées en cuivre ou en fibre optique, le choix dépendant de la distance.

L'agrégation de liens offre ici une protection supplémentaire car chaque lien passe par deux gaines techniques différentes : en cas de perte d'une des liaisons (coupure du câble, incendie, etc.), l'autre reste opérationnelle et permet un fonctionnement en mode dégradé à 1 Gbit/s au lieu de 2 Gbit/s.

Il est à noter que cette architecture redondante n'est possible que parce que les équipements sont des commutateurs. Ceux-ci échangent des informations sur la topologie du réseau et déterminent les routes actives et celles qui ne le sont pas. Le protocole utilisé s'appelle le **spanning tree**.

La même architecture avec des concentrateurs serait impossible car le segment Ethernet ferait une boucle, ce qui est interdit par la norme IEEE. Pour assurer le même niveau de redondance, une configuration équivalente avec des concentrateurs serait donc beaucoup plus complexe et nécessiterait de doubler le nombre de câbles. De plus, seule l'utilisation de la fibre optique permet ce type de redondance.

Figure 6-10.
*Redondance
avec des concentrateurs
Ethernet.*



Les liaisons en pointillé sont inactives en fonctionnement normal. En cas de rupture de la liaison principale, le concentrateur d'étage active la seconde liaison en quelques microsecondes.

La liaison entre les deux concentrateurs fédérateurs doit être doublée. En effet, si elle était perdue, on se retrouverait avec deux réseaux isolés ayant le même numéro de réseau IP.

QU'EST-CE QUE LE SPANNING TREE ?

Lorsqu'un commutateur reçoit une trame, il recherche son adresse de destination dans ses tables d'adresses MAC (une par port). Si elle ne s'y trouve pas, il la transmet sur tous ses ports. Sinon, il l'envoie uniquement sur le port identifié. C'est le mécanisme de **forwarding**.

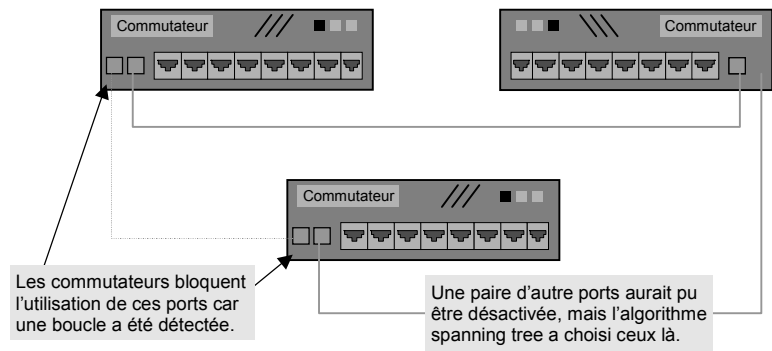
S'il existe plusieurs chemins, par exemple entre un commutateur d'étage et deux fédérateurs, la trame risque de boucler, car chacun des commutateurs transmet la trame (*forward*).

Le spanning tree est un **protocole de routage de niveau 2** associé à un algorithme qui permet d'éviter que les trames bouclent dans le réseau. Les commutateurs s'échangent des trames *spanning tree* et calculent une route en invalidant les chemins multiples susceptibles de créer des boucles au sein du segment Ethernet.

Inversement, avec le spanning tree, les liens sont tous actifs, mais une des routes est invalidée par les commutateurs. En cas de rupture d'un lien, une nouvelle route sera calculée, généralement au bout de trente secondes.

Figure 6-11.

La redondance avec le spanning tree.



Le fonctionnement du spanning tree est détaillé au chapitre 7.

DEUXIÈME PARTIE

Interconnecter ses réseaux

Démarrer son réseau IP

Des architectures viennent d'être élaborées, des réseaux viennent d'être construits, des protocoles sont utilisés. Mais comment tout cela s'imbrique-t-il ? Comment IP, TCP et Ethernet fonctionnent-ils ensemble ?

Ce chapitre est l'occasion d'aller plus loin dans la connaissance de ces protocoles et de votre réseau, et de vous donner par la même occasion une vision plus globale des réseaux. Car, comprendre, c'est pouvoir construire des réseaux de plus en plus complexes comme le requièrent les applications multimédias d'aujourd'hui.

Jusqu'à présent, nous ne nous sommes préoccupés que du matériel mais, avec le spanning tree, introduit au chapitre précédent, nous devons désormais nous préoccuper du paramétrage logiciel des équipements réseau.

Comprendre, c'est donc maîtriser le fonctionnement de son réseau.

Deuxième exemple, celui de l'adresse IP que nous avons utilisé sans en bien comprendre les tenants et aboutissants. Cet aspect logiciel doit maintenant être expliqué, car les choix que vous prenez lorsque vous commencez par construire un petit réseau peuvent ensuite peser bien des années plus tard, lorsque celui-ci a pris de l'ampleur.

Comprendre, c'est donc anticiper et faire les bons choix pour l'avenir.

Dans ce chapitre, vous apprendrez ainsi :

- à définir un plan d'adressage IP ;
- à comprendre et à paramétrer le spanning tree ;
- le fonctionnement d'un réseau local ;
- le fonctionnement des protocoles IP, TCP et UDP.

Le plan d'adressage IP

À plusieurs reprises déjà, nous avons parlé d'adresses IP sans vraiment nous en préoccuper. Il est vrai que nous n'en avons pas réellement l'usage ; il suffisait simplement de saisir une adresse unique pour chaque station du réseau.

Mais notre réseau prend de l'ampleur et nous devons désormais organiser l'affectation des adresses. Les ISP ne procèdent pas autrement au sein de l'Internet.

Comme cela a été dit au chapitre 4, une adresse IP s'écrit avec quatre numéros, compris entre 1 et 255, séparés par des points, par exemple 192.162.0.1. Une partie de cette adresse désigne un réseau, l'autre le numéro de station au sein de ce réseau. Jusqu'à présent, nous nous sommes arrangés pour configurer toutes nos stations dans le même réseau IP.

On peut se poser la question suivante : pourquoi faut-il des adresses IP alors qu'il existe déjà des adresses MAC ? D'abord, Ethernet est un réseau local, qui n'a donc qu'une portée géographique limitée. Ensuite, il existe des dizaines de réseaux de niveau 1 et 2 différents avec chacun un adressage physique qui lui est propre. Or, les PC, même situés sur des réseaux différents, doivent pouvoir communiquer ensemble. Il faut donc un protocole de niveau supérieur, dit de niveau 3 (couche réseau), qui permet de fédérer ces réseaux avec un adressage unique. On trouve ainsi IP sur Ethernet et PPP, mais aussi sur Token-Ring, ATM, etc.

IP permet aussi de partitionner les réseaux. En effet, de nombreux protocoles utilisent abondamment les broadcasts et multicasts, et il est préférable de limiter la diffusion de ces types de trames. Si votre intranet est connecté à l'Internet, il n'est pas envisageable de recevoir des trames multicast et broadcast émises par un employé de la société X.

De plus, l'interconnexion des sites coûte cher compte tenu des distances. Il est donc judicieux de limiter le trafic afin de ne pas surcharger inutilement les liaisons par des broadcasts.

POURQUOI UN PLAN D'ADRESSAGE ?

L'objectif premier du plan d'adressage est d'éviter la **duplication** accidentelle des adresses. Pour l'adressage MAC, un plan n'est pas utile car les adresses sont affectées aux cartes par les constructeurs. En revanche, l'affectation des adresses IP relève de votre responsabilité, ou de celle du NIC pour le réseau public Internet.

Le plan d'adressage permet également de **contrôler** le fonctionnement de votre réseau IP. En effet, l'affectation des adresses IP doit répondre à des règles précises sous peine d'aboutir à des dysfonctionnements (connexions impossibles, voire intermittentes, etc.).

En définitive, le plan d'adressage permet **d'organiser l'exploitation** de votre intranet.

La démarche

Tout d'abord, il est conseillé de retenir un adressage privé, c'est-à-dire complètement séparé de celui de l'Internet, ceci pour des questions de simplicité et de sécurité. Il est toujours possible d'opter pour un adressage public, mais l'obtention de telles adresses est très difficile car il faut justifier de leur usage auprès des organismes de régulation de l'Internet.

L'ADRESSAGE IP (RFC 791)

IP (*Internet Protocol*) définit un **réseau virtuel** reposant sur des réseaux physiques de différente nature (Ethernet et PPP, par exemple). Pour ce faire, IP utilise un **adressage logique** différent de l'adressage physique (MAC, PPP ou autre).

Une adresse IP est découpée en un numéro de réseau et un numéro de station au sein de ce réseau. Il existe trois **classes d'adresses** unicast en fonction de la taille du réseau (c'est-à-dire du nombre de stations par réseau). Pour différencier la partie réseau (*subnet*) de la partie station (*host*), IP utilise un **masque** dont tous les bits à 1 représentent la partie réseau.

Classe A		Masque naturel = 255.0.0.0	
0	7 bits pour le n° de réseau de 1 à 127	24 bits pour le n° de station, de 1 à 16 777 214	
126 réseaux de 1.0.0.0 à 126.0.0.0		0.x.x.x	Réservé
		127.x.x.x	Adresse de boucle locale (loopback)
		x.255.255.255	Broadcast : toutes les stations sur le réseau x
Classe B		Masque naturel = 255.255.0.0	
1	0	14 bits pour le n° de réseau, de 1 à 16 383	16 bits pour le n° de station, de 1 à 65 534
16 382 réseaux de 128.1.0.0 à 191.254.0.0		128.0.x.x	Réservé
		191.255.x.x	Réservé
		x.x.255.255	Broadcast : toutes les stations sur le réseau x.x
Classe C		Masque naturel = 255.255.255.0	
1	1	0	21 bits pour le n° de réseau, de 1 à 2 097 151
2 097 150 réseaux de 192.0.1.0 à 223.255.254.0		192.0.0.x	Réservé
		223.255.255.x	Réservé
		x.x.x.255	Broadcast : toutes les stations sur le réseau x.x.x

Deux valeurs sont réservées dans la partie station de l'adresse : 0 pour désigner le réseau lui-même et 255 (tous les bits à 1) pour désigner toutes les stations au sein de ce réseau (broadcast).

Il existe également une classe d'adresses multicast permettant de désigner des groupes de stations.

Classe D		Pas de masque	
1	1	1	0
268 435 455 groupes de 224.x.x.x à 239.255.255.255		28 bits pour le n° de groupe, de 1 à 268 435 456	
		224.0.0.0	Réservé
		224.0.0.1	Tous les groupes sur ce réseau local
		Des n° sont déjà réservés (<i>well known group</i>)	

La classe E (premiers à bits 11110) définit une classe d'adresses expérimentales. Elle n'est jamais utilisée.

L'adresse 255.255.255.255 désigne toutes les stations sur le réseau de l'émetteur du paquet (broadcast IP).

Il se peut donc que vous utilisiez des adresses déjà affectées sur l'Internet, mais cela n'a pas d'importance car votre intranet est isolé. Cela ne vous empêchera cependant pas de l'interconnecter avec l'Internet.

La seconde décision concerne le choix de la classe d'adresse IP. Ce choix dépend du nombre de stations présentes sur votre réseau. Si ce nombre dépasse 254, une classe B s'impose. Une classe A n'est pas utile, car une classe B offre 65 534 adresses de stations, ce qui est largement suffisant. De plus, une classe A est limitée à 126 réseaux IP, ce qui, pour les grands réseaux, peut être un handicap.

En résumé, notre choix s'est provisoirement porté sur un plan d'adressage privé de classe B, ce qui nous donne 16 382 réseaux possibles contenant chacun 65 534 stations. Aux sections suivantes, d'autres considérations viendront modifier ce choix.

Les principes de base

L'adressage IP est très souple et permet de faire tout ce que l'on veut. Afin d'éviter toute mauvaise surprise, il est conseillé de suivre les principes suivants :

- Règle 1 : un réseau IP ne doit pas chevaucher plusieurs sites.
- Règle 2 : il peut y avoir plusieurs réseaux IP sur un site.
- Règle 3 : s'il y a plusieurs réseaux IP sur un site, choisir des numéros contigus. Cela simplifiera le routage.
- Règle 4 : limiter le nombre de réseaux IP. Cela simplifiera les connexions à l'Internet.

Le protocole IP impose qu'une station se trouvant dans un réseau IP ne puisse pas communiquer directement avec une station se trouvant dans un autre réseau IP, même si elles sont connectées au même segment Ethernet. Les réseaux sont segmentés de manière logique ; en d'autres termes, ils sont partitionnés.

La solution repose sur l'utilisation d'un **routeur** dont le rôle est d'interconnecter les réseaux IP, quelle que soit leur localisation géographique.

On verra au chapitre 11 qu'il existe un moyen de lever cette contrainte imposée par IP.

De toute façon, l'utilisation d'un routeur s'impose dès que vous devez relier deux sites sur de longues distances. L'Internet comporte des dizaines de milliers de routeurs. Donc, autant prendre en compte cette contrainte dès le début de l'élaboration du plan d'adressage.

QU'EST-CE QU'UN ROUTEUR ?

Un routeur est un commutateur de niveau 3, c'est-à-dire qui commute les protocoles de la couche réseau, tels que IP. La commutation des paquets IP est plus complexe que celle des trames Ethernet. On emploiera donc plutôt le terme de **routage**.

Ce mécanisme consiste à analyser l'adresse de destination du paquet IP et à le transmettre sur le bon port (appelé **interface**). Il utilise pour cela des algorithmes de routage, tels que **OSPF** (*Open Shortest Path First*) qui permettent de calculer les meilleures routes en fonction des numéros de réseau IP.

Comme pour les PC, une interface routeur est associée à au moins un réseau IP.

Impact sur l'Internet

Tôt ou tard, l'interconnexion de votre réseau avec l'Internet sera nécessaire. Comment éviter que vos adresses internes entrent en conflit avec celles de l'Internet ?

La solution repose sur l'utilisation de la **translation d'adresses**. Cette technique permet de masquer votre plan d'adressage privé vis-à-vis des utilisateurs situés sur l'Internet.

Une solution complémentaire à la première repose sur le non-routage de certaines adresses. La RFC 1918 précise que certaines adresses ont été réservées pour l'adressage privé. Le respect par tous les ISP de cette RFC garantit que ces adresses ne seront jamais routées sur l'Internet.

Réseaux réservés (RFC 1918)	Espace d'adressage
10.0.0.0	1 réseau de classe A
De 172.16.0.0 à 172.31.0.0	16 réseaux de classe B
De 192.168.0.0 à 192.168.255.0	256 réseaux de classe C

On peut donc utiliser ces adresses pour notre réseau privé, sans que cela soit pour autant une obligation. L'essentiel d'une interconnexion avec l'Internet repose, en effet, sur la translation d'adresses.

Or, pour les grands réseaux, le nombre de réseaux IP à tradater est source de complexité : s'il y a quarante sites mais un seul point de sortie vers l'Internet, le firewall devra prendre en compte quarante réseaux IP dans ses règles de translation d'adresses.

Afin de simplifier cette configuration, il faudrait donc pouvoir ne tradater qu'un réseau IP au niveau du firewall (respect de la règle 4) tout en ayant autant de subnets IP que nécessaire pour notre intranet. La solution repose sur la création de sous-réseaux IP.

Les sous-réseaux IP

Le principe des sous-réseaux (*subnet*) consiste à étendre le nombre de bits désignant la partie réseau. Le nombre de stations par sous-réseau diminue donc d'autant.

Classe	Masque naturel	Nombre de bits affectés au numéro de réseau	Extension possible : nombre de bits affectés au sous-réseau
A	255.0.0.0	8	+ 1 à + 22 bits
B	255.255.0.0	16	+ 1 à + 14 bits
C	255.255.255.0	24	+ 1 à + 6 bits

La partie station de l'adresse doit comporter au moins 2 bits afin que cette dernière soit valide.

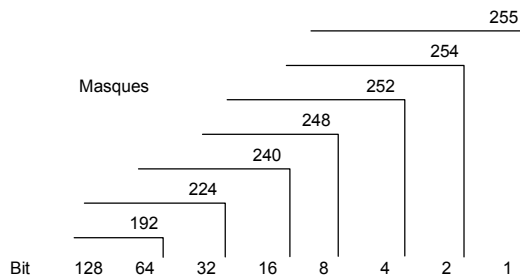
	Numéro de station	Conclusion
1 bit	1 = broadcast	Interdit. Il ne reste aucun bit pour désigner un sous-réseau ou une station.
	0 = ce réseau	
2 bits	11 = broadcast	OK. 2 est le nombre minimal de bits devant être réservés aux sous-réseaux et stations.
	00 = ce réseau	
	01 = station n° 1	
	10 = station n° 2	

La notation décimale (octet par octet) est rendue difficile lorsque le sous-réseau ne porte pas sur un multiple de 8 bits. C'est pourquoi la notation « / [nombre de bits affectés à la partie réseau] » est plus souvent utilisée.

Numéro de réseau / nombre de bits réservés à la partie réseau	Masque	Commentaire
10.0.0.0 / 10 subnet de +2 bits	255.192.0.0	Permet de créer 4 sous-réseaux, de 10.0 à 10.3
10.0.0.0 / 16 subnet de +8 bits	255.255.0.0	On dit que la classe A est « subnettée » sur une classe B
194.50.0.0 / 19 subnet de +3 bits	255.255.224.0	Permet de créer 8 sous-réseaux
194.50.0.0 / 24 subnet de +8 bits	255.255.255.0	On dit que la classe B est « subnettée » sur une classe C

Figure 7-1.

Les masques de sous-réseaux.



Notre choix initial portait sur une classe B. Si nous voulons limiter le nombre de réseaux IP et conserver la même souplesse que la classe B, il faut donc retenir une classe A « subnettée » sur une classe B.

Cela nous offrirait 256 sous-réseaux. Si, dans le futur, ce chiffre était dépassé, on pourrait toujours ajouter un autre réseau de classe A (il ne ferait pas partie de la RFC 1918, mais ce n'est pas réellement important) et le « subnetter », ou ajouter une classe B à notre plan d'adressage. Notre but est simplement de limiter le nombre de réseaux IP.

Nous choisissons donc l'adresse de classe A, 10.0.0.0, issue de la RFC 1918. Étant donné le subnet choisi, notre masque sera donc : 255.255.0.0. Mais ce choix est encore provisoire.

Méthode d'affectation des réseaux LAN

Le plus simple est d'affecter les réseaux par site (respect de la règle 1). Au lieu d'affecter séquentiellement le numéro, on peut l'incrémenter de 4 ou 8, ce qui laisse la possibilité d'étendre le subnet affecté au site (respect de la règle 2). L'ajout d'un réseau sur un site se traduira donc par l'affectation du numéro de réseau suivant (respect de la règle 3).

Une première version de notre plan d'adressage serait donc la suivante :

Réseau	Site
10.0.0.0/16	Paris : 1 réseau de 65 534 stations
De 10.1.0.0/16 à 10.3.0.0/16	Non affecté (réservé aux extensions de Paris)
10.4.0.0/16	Toulouse : 1 réseau de 65 534 stations
De 10.5.0.0/16 à 10.7.0.0/16	Non affecté (réservé aux extensions de Toulouse)
Etc.	
De 10.248.0.0 à 10.255.0.0	Réseaux non affectés

L'incrément de 4 a été soigneusement choisi, de manière à obtenir des réseaux contigus. Ainsi, le site de Paris dispose de quatre réseaux : 10.0.0.0, 10.1.0.0, 10.2.0.0 et 10.3.0.0, avec chacun un masque à 255.255.0.0. Mais cette manière de découper les réseaux est quelque peu rigide, car la région de Paris peut comprendre à la fois des petits sites et des gros sites.

Une autre façon de voir les choses est de considérer le réseau 10.0.0.0 avec le masque 255.252.0.0 (soit 10.0.0.0/14), ce qui offre 262 142 adresses ($65\,536 \times 4 - 2$) pour le subnet 10.0.0.0 affecté à Paris (de 10.0.0.0 à 10.3.255.255).

Figure 7-2.
Création d'un subnet.

Adresse	10.						0.		0.0
Masque	255.						252.		0.0
Bit	128	64	32	16	8	4	2	1	
/14	8 bits						+ 6 bits		

En définitive, notre plan d’adressage se présente en réalité sous la forme suivante :

Subnets du réseau 10.0.0.0/8	Site
10.0.0.0/14 255.252.0.0	Région parisienne
10.4.0.0/14 255.252.0.0	Région toulousaine
10.8.0.0/14 255.252.0.0	Strasbourg
Etc.	
De 10.248.0.0 à 10.255.252.0	Réseaux non affectés

Au sein de ce réseau, il est alors possible de créer d’autres subnets dont la taille varie en fonction de l’importance du site. En faisant varier la longueur du masque on crée ainsi des **subnets variables** (RFC 1219).

Par exemple, au sein de la plage d’adresses affectée à la région parisienne, on peut réserver le subnet suivant à un site de moyenne importance : 10.0.0.0/20 (masque égal à 255.255.240.0), soit 4 094 adresses ($16 \times 256 - 2$), de 10.0.0.1 à 10.15.255.254.

Figure 7-3.
Création d’un deuxième subnet.

Adresse	10.0.	0.	0					
Masque	255.255.	240.	0					
Bit	128	64	32	16	8	4	2	1
/20	16 bits			+ 4 bits				

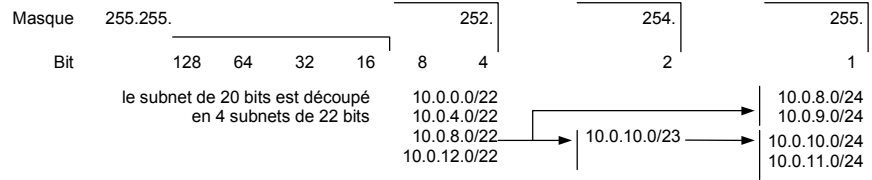
Au sein de ce site, il peut ensuite être nécessaire de créer des réseaux de différentes tailles, par exemple un réseau principal et de nombreux petits sous-réseaux dédiés connectés, par exemple, à un firewall.

Subnets du réseau 10.0.0.0/14	Fonction
10.0.0.0/22 255.255.252.0	Réseau principal (1 022 adresses)
10.0.4.0/22 255.255.252.0	Réservé à l’extension du réseau principal (*) ou à la création d’un deuxième réseau
10.0.8.0/24 255.255.255.0	Réseaux dédiés au firewall (254 adresses)

(*) Si le réseau principal est étendu, il suffit de changer le masque qui devient 255.255.248.0, ce qui donne le réseau 10.0.0.0/21.

Le réseau 10.0.10.0/23 peut également être découpé en deux subnets de classe C (masque de 24 bits) 10.0.10.0 et 10.0.11.0.

Figure 7-4
*Extension
des subnets.*



Les subnets de classe C ainsi créés (10.0.8.0, 10.0.9.0, etc.) peuvent à leur tour être découpés en de très petits réseaux, juste assez grands pour connecter un routeur et quelques machines.

Subnets du réseau 10.0.8.0/24	Fonction du réseau dédié
10.0.8.0/27 255.255.255.224	Serveurs publics (30 adresses)
10.0.8.32/27 255.255.255.224	Réservé (30 adresses)
10.0.8.64/26 255.255.255.192	Accès distants (62 adresses)
10.0.8.128/28 255.255.255.240	Accès externes (14 adresses)
10.0.8.144/28 255.255.255.240	Réservé (14 adresses)
10.0.8.160/27 255.255.255.224	PABX (30 adresses)
10.0.8.192/26 255.255.255.192	Réservé (62 adresses)

Une autre manière d'appréhender la subtilité du subnetting qui vient d'être opéré est de considérer la grille de découpage suivante.

	2 subnets de 128 (– 2) adresses	4 subnets de 64 (– 2) adresses	8 subnets de 32 (– 2) adresses	16 subnets de 16 (– 2) adresses
Plage d'adresses au sein du subnet	0 – 127	0 – 63	0 – 31	0 – 15
			32 – 63	16 – 31
		64 – 127	64 – 95	32 – 47
			96 – 127	48 – 63
				64 – 79
				80 – 95
	128 – 255	128 – 192	126 – 159	96 – 111
			160 – 191	112 – 127
		192 – 255	192 – 223	128 – 143
			224 – 240	144 – 159
				160 – 175
				176 – 191
				192 – 207
				208 – 223
				224 – 239
				240 – 255
Masque	/25	/26	/27	/28

Les plages réservées permettront d'étendre les plages déjà affectées si le nombre de stations devient plus important que prévu. Ainsi, le réseau " Serveurs publics " pourra être étendu en diminuant le masque de 1 bit, afin de donner le subnet 10.0.8.0/26 (255.255.255.192).

Il est à noter que la création d'un sous-réseau fait perdre chaque fois deux adresses.

La technique du *subnetting* permet de gérer la pénurie d'adresses publiques sur l'Internet. En effet, la création de réseaux IP taillés sur mesure évite le gaspillage d'adresses ; par exemple, le réseau 10.0.0.0/16 offre 65 534 adresses qui seront loin d'être toutes utilisées. Sur votre réseau privé, vous avez cependant plus de latitude. Mais attention aux évolutions qui peuvent être rapides, par exemple lors de la fusion de deux sociétés.

Méthode d'affectation des réseaux WAN

L'interconnexion des réseaux (abordée aux chapitres suivants) nécessite également des adresses, mais en moins grand nombre que pour les réseaux LAN.

Par exemple, sur une liaison point à point, seules deux adresses sont nécessaires, une pour chaque extrémité. Le subnetting sur 30 bits, qui offre deux adresses, permet de créer un réseau juste dimensionné pour ce besoin.

Nous pourrions utiliser une des plages de notre réseau 10, mais il est cependant plus intéressant d'utiliser un autre réseau IP, et cela pour plusieurs raisons :

- Les adresses des réseaux WAN ne sont pas diffusées sur l'ensemble du réseau ; elles ne sont connues qu'entre routeurs adjacents.
- Les adresses n'ont donc pas besoin d'être connues des réseaux utilisateurs.
- Utiliser une plage d'adresses distincte permet de mieux identifier les liaisons WAN.

Bien que cela ne soit pas une obligation, nous préférons donc utiliser une autre plage d'adresses de la RFC 1918. Une classe B suffira amplement.

Nous pouvons donc réserver une plage de notre réseau 172.16.0.0/16, que nous « subnetterons » comme suit :

Subnets de 172.16.0.0/16	Fonction
172.16.0.0/30 255.255.255.252	Liaison Paris-Toulouse
172.16.0.4/30 255.255.255.252	Liaison Paris-Strasbourg
etc.	En tout : 16 384 subnets de 2 adresses

Pour les interconnexions multipoints, il suffira de réduire le masque d'autant de bits que nécessaire pour les subnets considérés. En général, les réseaux multipoints WAN sont rares et comprennent peu d'adresses en comparaison des LAN.

Méthode d'affectation des stations au sein des réseaux

Chaque nœud IP doit posséder une adresse IP. Cela concerne les PC et les Macintosh, les serveurs (NT, Unix, etc.), les imprimantes, les routeurs, les concentrateurs et commutateurs administrables (pour les agents SNMP), etc.

Il est tentant de découper la plage d'adresses en autant de parties qu'il y a de types de matériels. Cela n'apporterait cependant rien ni sur un plan technique, ni sur un plan organisationnel.

L'expérience montre, de plus, qu'une telle pratique n'est pas pérenne : soit la taille de la plage que l'on avait réservée est insuffisante (davantage de PC que prévu, par exemple), soit, à la longue, personne ne respecte une marche à suivre qui est trop contraignante (par exemple, s'il faut installer un PC en urgence, on prend la première adresse disponible).

Il est, en revanche, intéressant de prévoir un découpage simple entre les équipements terminaux (PC, serveurs, imprimantes, etc.) et les équipements réseau (les routeurs, les agents SNMP des concentrateurs et des commutateurs, etc.). Cela permet de mieux contrôler les flux du réseau. Dans le cas d'un subnet de classe B, on peut se risquer à créer une troisième plage réservée aux serveurs.

Plage d'adresses	Affectation
De 0.1 à 24.255	Équipements réseau (routeurs, hubs, switches, etc.). 6 399 adresses (de 1 à 6 399)
De 25.0 à 49.255	Serveurs NT, Unix, etc. 6 400 adresses (de 6 400 à 12 799)
De 50.0 à 255.254	Postes de travail (PC, etc.) 52 735 adresses (de 12 800 à 65 534)

Dans le cas d'un subnet sur une classe C, le plus simple est de ne pas affecter de plage d'adresses par type d'équipement, car la probabilité de collision est encore plus forte qu'avec une classe B. L'affectation des adresses pour les équipements réseau et serveur peut commencer par le bas de la plage et s'incrémenter ensuite, tandis que celle pour les PC peut commencer par le haut de la plage et se décrémenter ensuite.

Plage d'adresses	Affectation
De .1 à .254	Équipements réseau (routeurs, hubs, switches, etc.) et serveurs NT, Unix, etc.
De .254 à .1	Stations de travail (PC, etc.)

On peut constater que le plan d'adressage doit prendre en compte de nombreux paramètres liés à des notions qui n'ont pas été introduites : routage, translation d'adresse, affectation dynamique, connexion à l'Internet et contrôle de flux. Les chapitres suivants vous permettront de juger de la pertinence ou non du plan d'adressage qui vous est proposé.

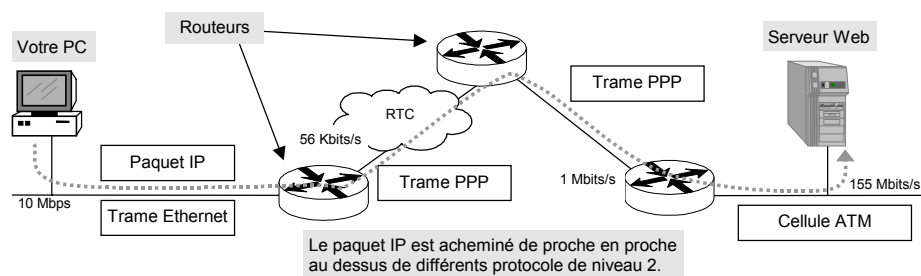
L'encapsulation des protocoles

Il existe différents supports de transmission (câbles en cuivre ou en fibre optique, faisceaux hertziens) et différents moyens d'accéder à ces supports (accès partagé par détection de collision, par jeton, par partage fixe de bande passante, etc.). Cela implique l'utilisation de nombreux protocoles de niveau 1 (couche physique) adaptés à chaque situation.

La couche liaison, telle que PPP, permet de masquer aux couches supérieures les particularités du niveau physique et ses contraintes. Mais il arrive qu'une norme spécifie les couches 1 et 2 : c'est le cas d'Ethernet et d'ATM (*Asynchronous Transfert Mode*).

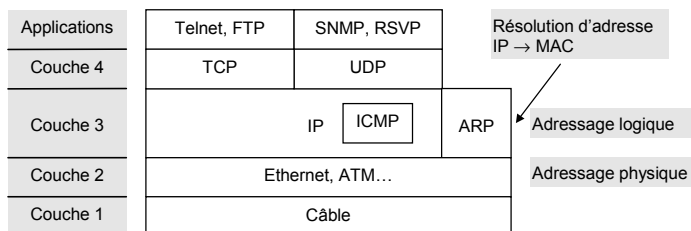
La couche de niveau 3 (couche réseau), telle que IP, peut donc utiliser différents réseaux en recourant aux services de PPP ou en s'adaptant directement sur une autre couche liaison.

Figure 7-5.
Paquet IP
au-dessus
de différents
réseaux.



On peut établir l'analogie suivante : le paquet IP est une voiture ; les pneus et les suspensions sont les protocoles de niveau 2 qui réalisent l'adaptation aux routes que sont les réseaux physiques. Vous roulez ainsi sur un chemin de terre (le RTC), puis sur une nationale (Ethernet) et enfin sur une autoroute (ATM), mais toujours avec la même voiture. Éventuellement, vous changez de pneus ou de suspensions, afin de vous adapter au terrain. De même, le paquet IP peut emprunter le RTC (avec une trame PPP), un réseau Ethernet (avec une trame Ethernet) ou un réseau ATM (avec une cellule ATM).

Figure 7-6.
Modèle en couches
des protocoles Internet.



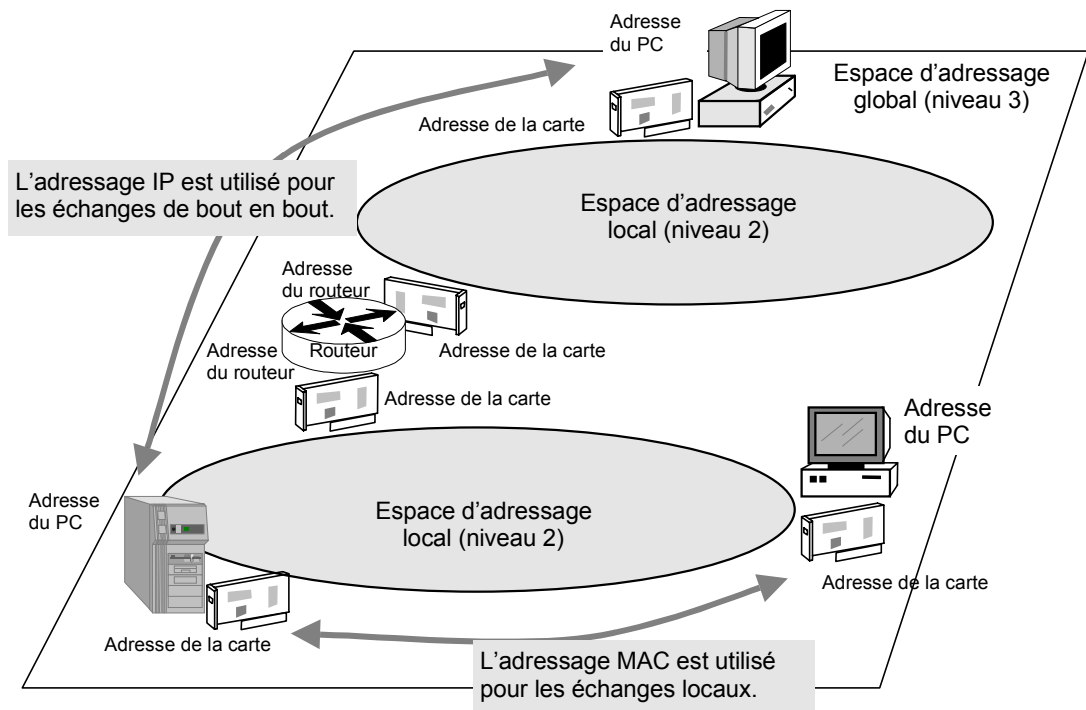
L'adressage

Toutes les couches réseau, de la couche physique à l'application en passant par les couches liaison, réseau et transport, utilisent des adresses afin d'identifier l'émetteur et le destinataire. Chaque couche utilise un système d'adressage spécifique qui répond à un besoin précis.

L'adressage de niveau 2 est géographiquement limité à un réseau local ou à une liaison point à point d'un réseau étendu.

L'adressage de la couche 3 permet d'identifier les stations à un niveau supérieur. Il assure la continuité entre des réseaux physiques qui utilisent différents systèmes d'adressage.

Figure 7-7.
Le rôle de l'adressage.



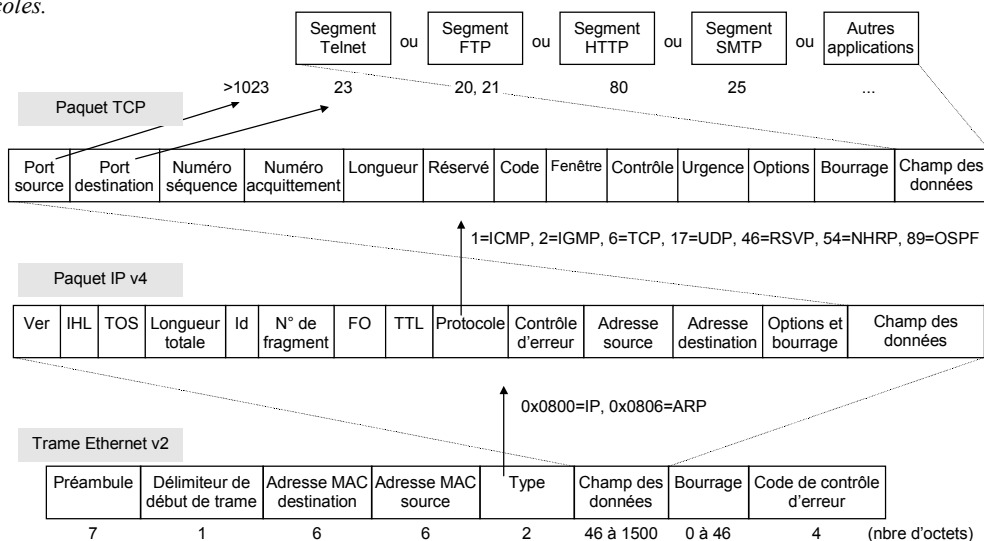
Enfin l'adressage de niveau 4 permet d'identifier les applications qu'utilisent les services de la couche transport.

Le multiplexage

Chaque couche réseau dispose d'un champ pour identifier le type de protocole encapsulé dans le champ de données. Ethernet identifie ainsi qu'il transporte un paquet IP, IP identifie qu'il transporte des données TCP, et TCP identifie l'application qui a rempli son champ de données.

Figure 7-8.

Le principe de l'encapsulation des protocoles.



Les champs “Type”, “Protocole” et “Port” permettent à chaque couche de savoir à quelle couche supérieure remettre les données reçues. La RFC 1700 recense ainsi toutes les valeurs affectées aux protocoles de la famille TCP/IP ou à ceux qui utilisent IP.

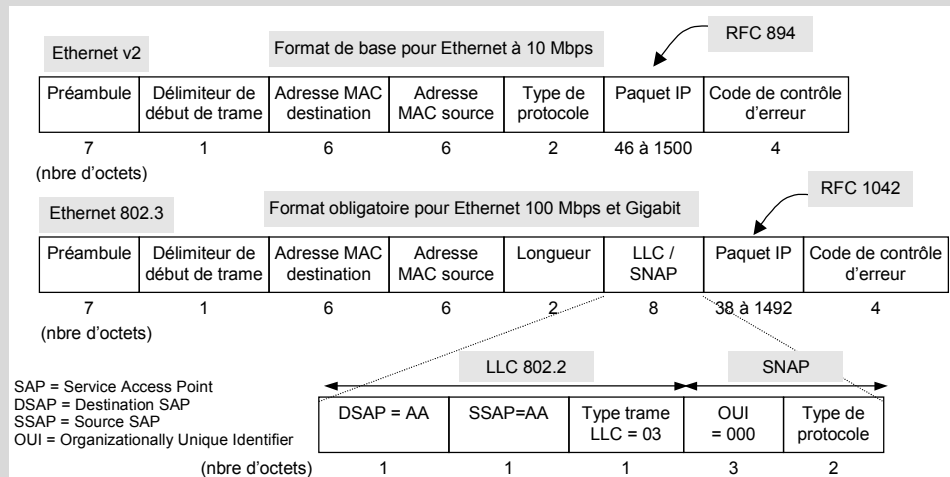
Il est ainsi possible d'envisager toutes les combinaisons d'encapsulation, telles que celle spécifiée par le protocole STUN (*Serial Tunneling*) qui permet de transporter dans un paquet IP une trame SDLC (*Synchronous Data Link Control*) qui est un protocole de niveau 2 utilisé dans les réseaux SNA d'IBM. On pourra ainsi trouver l'encapsulation : “SDLC → TCP → IP → SNAP → LLC → Ethernet 802.3”.

En théorie, tout protocole peut donc être encapsulé dans n'importe quel autre protocole. Dans la pratique, on utilise cette facilité pour répondre à une contrainte particulière, telle que le transport des flux SNA dans un réseau IP.

L'ENCAPSULATION D'IP DANS ETHERNET (RFC 894 ET 1042)

Il existe deux formats de trames : **Ethernet v2**, également appelée DIX (du nom des constructeurs Digital, Internet et Xerox), et **Ethernet IEE 802.3**. Il y a donc deux façons d'envoyer un paquet IP sur Ethernet : directement dans une trame Ethernet v2 (RFC 894) ou *via* un en-tête LLC/SNAP dans une trame 802.3 (RFC 1042).

Si la valeur du champ "Type de protocole/Longueur" est supérieure à 1 500 (correspondant au nombre maximal d'octets pour le champ contenant le paquet IP), il s'agit d'une trame Ethernet v2 (le champ a alors la signification "Type de protocole"). Dans le cas contraire, il s'agit d'une trame Ethernet 802.3 (le champ a alors la signification "Longueur").



Il existe différents types de **trames LLC** impliquant différents modes de fonctionnement. Pour IP, seule la trame de type *Unnumbered Information* (type 03) est utilisée (trame simple sans acquittement). Elle est également utilisée pour transporter IP dans ATM (*Classical IP*), Token-Ring et FDDI.

La **couche SNAP** (*Sub Network Access Protocol*) est nécessaire car la trame LLC (*Logical Link Control*) ne contient pas de champ équivalent au champ "Type" de la trame Ethernet v2. Le SAP (*Service Access Point*) utilisé pour transporter SNAP est 170 (0xAA). On retrouve ce principe d'adaptation avec d'autres protocoles comme Frame-Relay ou ATM (voir chapitre 10).

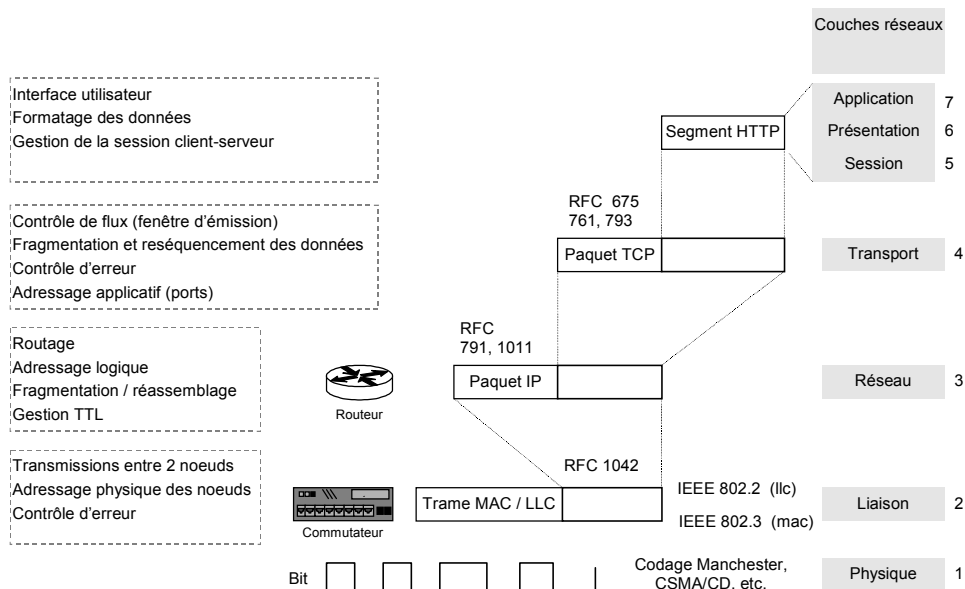
Dans le dernier champ de l'en-tête SNAP, on retrouve enfin le **type de protocole** utilisé dont les valeurs sont identiques à celles du champ de même nom de la trame Ethernet v2 (0x0800 pour IP, 0806 pour ARP, etc.).

Il est à noter que le **MTU** (*Maximum Transfert Unit*), c'est-à-dire les données utiles transportées dans la trame Ethernet, est plus important avec Ethernet v2, l'encapsulation 802.3 faisant perdre 8 octets.

Dans le cas d'une navigation sur le web, l'empilement des protocoles est : "segment HTTP → TCP (port 80) → IP (protocole 6)", puis toutes sortes de réseaux de transport, tels que PPP, Frame-Relay, ATM, etc.

Sur votre réseau, l'encapsulation sera : "IP → SNAP (protocole 2048) → LLC (SAP 170) → Ethernet 802.3".

Figure 7-9.
*Le rôle
des couches
réseau.*



Ce modèle en couches simplifie la programmation des protocoles en leur assignant des rôles précis, et offre plus de souplesse par le jeu des encapsulations possibles.

Comment une station envoie-t-elle une trame Ethernet à une autre ?

Un réseau local tel qu'Ethernet est constitué de concentrateurs et de commutateurs, deux équipements au comportement bien différent.

Caractéristique	Ethernet partagé	Ethernet commuté
Équipement	Concentrateur	Commutateur
Segment	Un segment partagé par tous les ports	Un segment par port
Architecture matérielle	Composants électroniques ; un ou plusieurs bus Ethernet	ASIC et processeur RISC, matrice de commutation
Architecture logicielle	Petit logiciel pour des options de configuration et pour l'administration SNMP	Logiciel plus complexe, mais traitement essentiellement matériel
Algorithme de routage	Aucun	Spanning tree
Traitement des trames	Aucun traitement (transparent pour les trames Ethernet)	Filtrage et transmission (<i>forward</i>) des trames ; apprentissage des adresses MAC
Couche réseau	Niveau 1 = couche physique	Niveaux 1 et 2 = couche physique et liaison
Fonction	Connecter plusieurs PC au sein d'un segment	Interconnecter plusieurs segments

Échange de trames sur un segment Ethernet

Commençons par étudier le fonctionnement au sein d'un segment Ethernet partagé (la trame circule sur un seul segment).

La norme Ethernet spécifie l'utilisation d'adresses physiques liées aux cartes réseau : les adresses MAC.

La carte recevant une trame Ethernet ne la prendra en compte que si l'adresse MAC de destination de la trame est identique à celle qui est inscrite dans sa mémoire. La seule exception à cette règle concerne les adresses de broadcast et, éventuellement, les adresses multicast. En résumé, les cartes sont programmées pour accepter les trames qui leur sont destinées, plus toutes les trames de broadcast ainsi que les trames multicast qui ont été configurées.

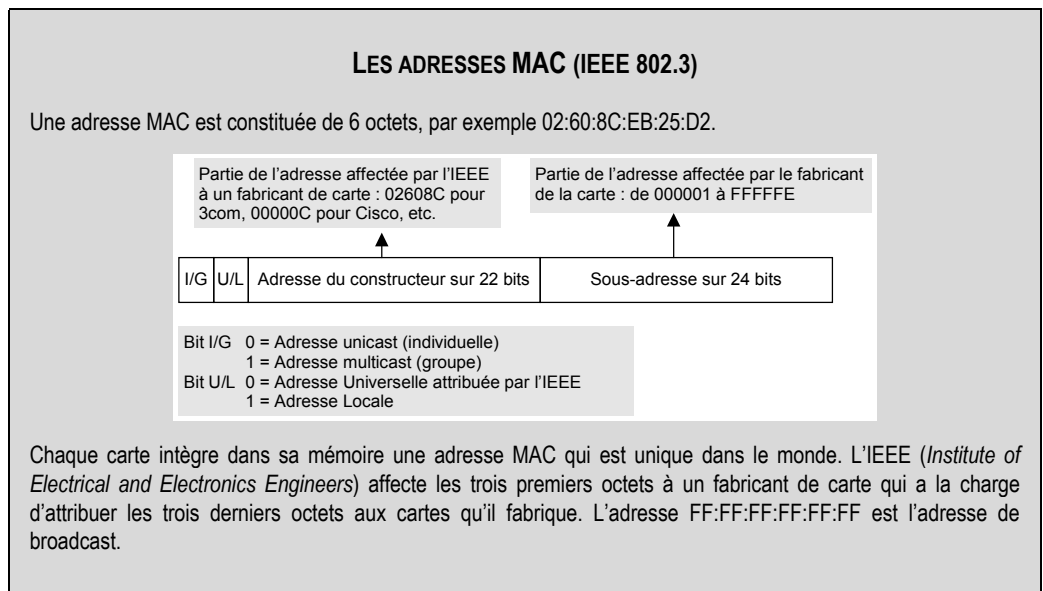
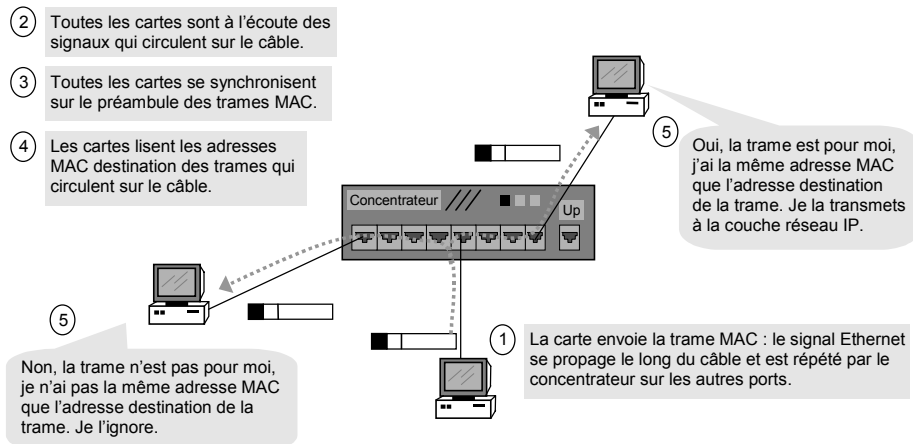


Figure 7-10.
Échange des trames Ethernet.

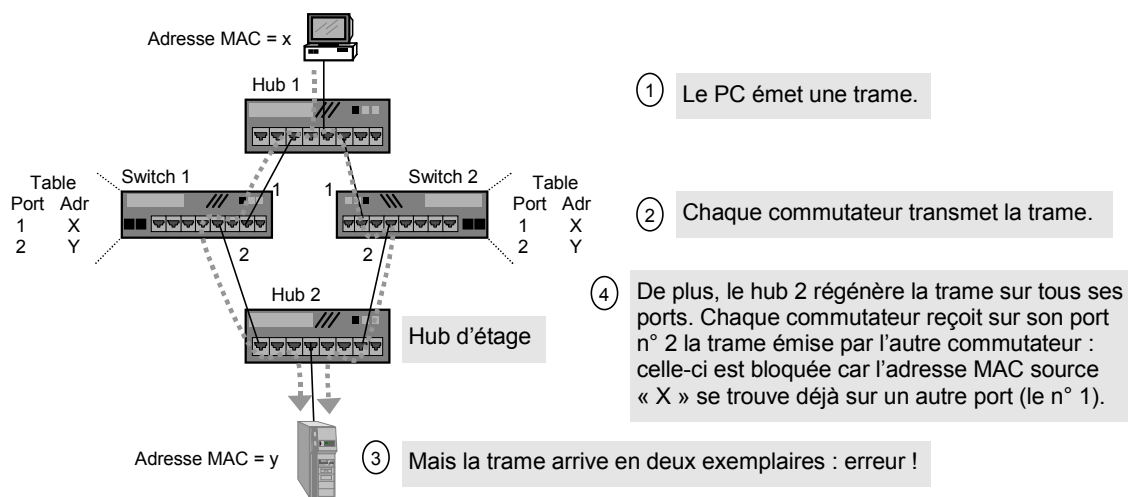


Les trames qui sont acceptées sont remises au protocole de niveau 3, qui correspond à l'identifiant trouvé dans le champ "Type", soit 0x0800 pour IP. Celles qui ne sont pas destinées à la station sont ignorées.

Échange de trames entre différents segments Ethernet

On l'a vu, le commutateur est un équipement permettant de segmenter le réseau Ethernet et d'interconnecter ces différents segments. Pour décider si la trame doit passer d'un segment à l'autre, il se base sur son adresse MAC qu'il compare avec celles se trouvant dans ses tables d'adresses en mémoire.

On peut envisager différentes situations dans lesquelles il existe plusieurs chemins possibles entre deux stations.



Il faut noter que, si le serveur n'a jamais émis de trame, son adresse MAC "Y" n'est pas encore connue des commutateurs. Ces derniers transmettront alors la trame sur tous les autres ports, dont le port n° 2. De plus, la même situation se produit pour toutes les trames de broadcast et de multicast.

Pour éviter ces problèmes, il faut qu'une des deux routes soit interdite. C'est là qu'intervient le **spanning tree**. Le but de ce protocole est de définir une route unique vers un commutateur désigné racine en se basant sur des coûts et des priorités.

LE POINT SUR LE SPANNING TREE (IEEE 802.1D)

L'algorithme du spanning tree consiste à construire un arbre définissant un chemin unique entre un commutateur et sa racine.

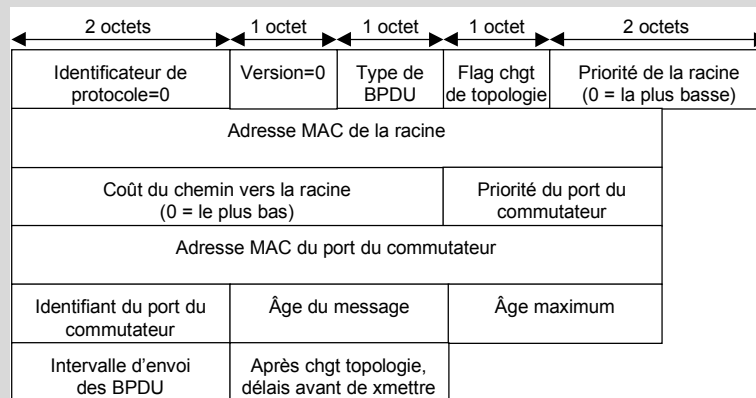
Lors de la construction de l'arbre (suite à un changement de topologie), chaque commutateur émet un **BPDU** (*Bridge Protocol Data Unit*) de configuration sur tous ses ports. Inversement, il retransmet tous les BPDU (éventuellement en les modifiant) qui lui arrivent, et ainsi de suite jusqu'à ce que les BPDU échangés contiennent tous la même valeur.

La première étape de ce processus consiste à élire un **commutateur racine** : c'est celui dont la **priorité** est la plus basse ou, en cas d'égalité, celui dont l'adresse MAC est la plus basse.

Ensuite, chaque commutateur détermine le **port racine** — celui par lequel un BPDU émis par la racine arrive. S'il y en a plusieurs, le port choisi est celui qui a le **coût** de chemin vers la racine le plus bas. Le coût est déterminé par la somme des coûts des ports situés entre le commutateur et la racine. En cas d'égalité, le port choisi est celui qui a la priorité la plus basse ; en cas de nouvelle égalité, c'est celui qui a l'adresse MAC la plus basse.

Enfin, sur chaque segment Ethernet, le commutateur dont le port racine a le coût de chemin vers la racine le plus bas est élu **commutateur désigné**. En cas d'égalité, c'est celui qui a la priorité la plus basse et, en cas de nouvelle égalité, celui qui a l'adresse MAC la plus basse.

En définitive, sur chaque segment Ethernet, un seul chemin vers le commutateur racine sera calculé. Les commutateurs désactivent tous leurs ports qui ne sont ni racines ni désignés.



Afin de détecter les changements de topologie (apparition ou disparition d'un commutateur), la racine envoie régulièrement (toutes les deux secondes par défaut) un BPDU d'annonce (comportant seulement les trois premiers octets) sur tous ses ports. Les commutateurs transmettent ce BPDU sur leurs ports désignés.

Les BPDU sont envoyés dans des trames Ethernet multicast 01:80:C2:00:00:10.

Compte tenu des processus d'élection, il est important de bien paramétrer les coûts (exprimés en nombre de sauts et/ou dépendant du débit du port) ainsi que les priorités. Les valeurs par défaut (fixées en usine) peuvent en effet aboutir à choisir des chemins qui ne sont pas les meilleurs. Prenons le cas de notre réseau local redondant.

Figure 7-11.

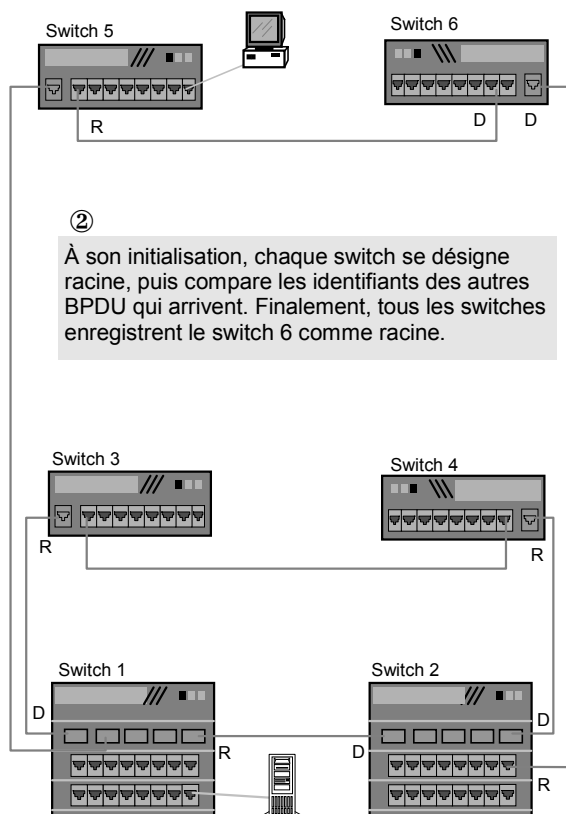
Conséquence d'un spanning tree mal paramétré.

① Tous les switches ont été installés avec les valeurs par défaut (valeurs usine). Ils ont les mêmes coûts et les mêmes priorités.

③

Résultat : le flux entre la station et le serveur transite par deux commutateurs au lieu d'être direct.

R = port racine
D = port désigné
En pointillé, les liens invalidés : les ports ont été bloqués par le spanning tree.



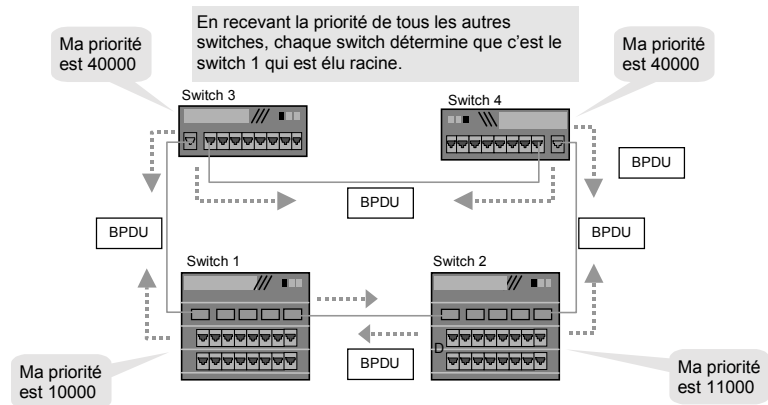
Tous les commutateurs ont la même priorité (par exemple 32 768) et le même coût sur chaque port de même débit (par exemple 19 pour les ports à 100 Mbit/s). Le switch 6 a été désigné racine parce que son identifiant (priorité + adresse MAC) était le plus bas. Le même processus de sélection a déterminé les routes menant vers la racine uniquement en se fondant sur les valeurs des adresses MAC, puisque toutes les autres valeurs (priorité et coût) sont identiques.

Résultat, certains flux ne sont pas optimisés et peuvent dégrader les performances.

Reprenons les différentes phases de calcul de l'arbre spanning tree. Les commutateurs désignent la racine. Afin d'optimiser les flux, il est préférable que ce soient les commutateurs fédérateurs qui assurent ce rôle. Leur priorité doit donc être abaissée par rapport aux commutateurs d'étage. Étant donné qu'il s'agit de Catalyst 5000, la commande est la suivante :

```
Console> (enable)set spanmtree priority 10000
VLAN 1 bridge priority set to 10000.
```

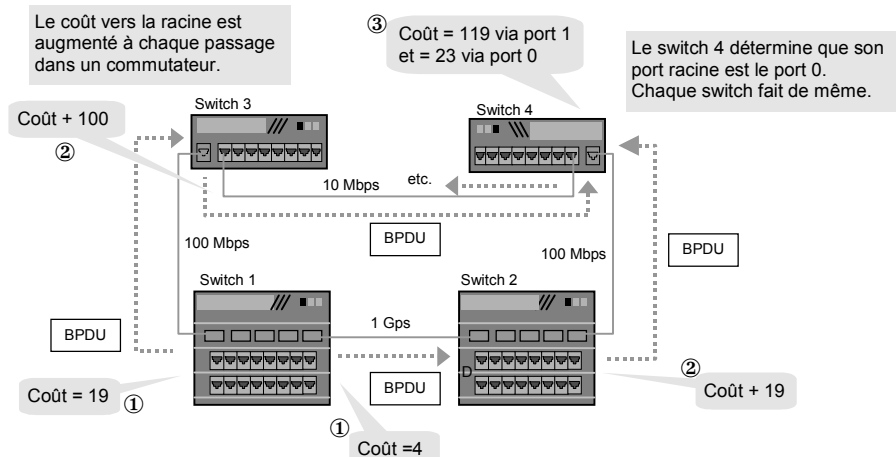
Figure 7-12.
*Élection
du commutateur racine.*



Chaque commutateur choisit ensuite son port racine, celui dont le coût de chemin vers la racine est le plus bas. Sur un Catalyst, le coût de chaque port dépend de son débit : 4 pour 1 Gbit/s, 19 pour 100 Mbit/s, et 100 pour 10 Mbit/s. La commande suivante permet de changer la valeur par défaut :

```
Console> (enable)set spantree portcost 1/1 4
```

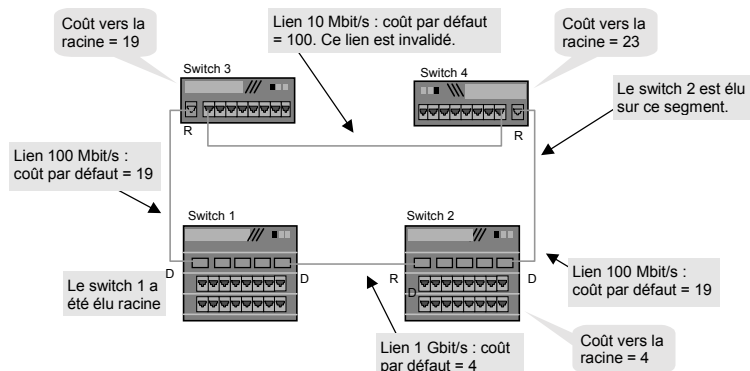
Figure 7-13.
*Choix
des ports racines.*



Sur chaque segment Ethernet, le commutateur désigné est celui dont le port racine a le coût le plus bas. En cas d'égalité, la priorité détermine ce coût. Étant donné que tous les ports ayant un même débit ont le même coût et la même priorité par défaut, le choix s'effectuera en fonction de l'adresse MAC. Pour éviter les mauvaises surprises, il est possible d'abaisser la priorité d'un port pour être sûr qu'il soit désigné en cas de routes multiples :

```
set spantree portpri 1/1 32
```

Figure 7-14.
*Élection
des commutateurs désignés.*



En définitive, les chemins redondants ne sont pas utilisés, et le partage de la charge entre plusieurs routes n'est pas possible.

Les mêmes BPDU sont envoyés sur tous les ports, même là où il n'y a qu'un PC connecté. Le spanning tree peut donc être désactivé sur ces ports, ce qui présente l'avantage de diminuer (un peu) le trafic et d'éviter que des ajouts sauvages de commutateurs (qui seraient connectés sur ces ports) ne viennent perturber votre réseau.

```
set spantree disable
```

Le commutateur racine émet régulièrement (toutes les deux secondes par défaut) des BPDU pour maintenir l'état du spanning tree. Si le réseau est stable (peu d'incidents et de changements), il est possible d'augmenter cette valeur afin de diminuer le trafic

```
set spantree hello 5
```

On peut s'assurer que, sur le switch 1, le spanning tree s'est stabilisé dans une bonne configuration.

```
Console> (enable) show spantree
VLAN 1
Spanning tree enabled
Designated Root          00-1f-00-40-0b-eb-25-d2
Designated Root Priority  45
Designated Root Cost     0
Designated Root Port     1/1
Root Max Age  20 sec    Hello Time 2  sec    Forward Delay 20 sec
Bridge ID MAC ADDR      00-40-0b-eb-25-d2
Bridge ID Priority      45
Bridge Max Age 20 sec    Hello Time 2  sec    Forward Delay 20 sec
```

Port	Vlan	Port-State	Cost	Priority	Fast-Start
-----	----	-----	-----	-----	-----
1/1	1	forwarding	4	32	disabled
1/2	1	forwarding	19	32	disabled
2/1	1	forwarding	19	32	disabled
2/2	1	not-connected	19	32	disabled
2/3	1	not-connected	19	32	disabled
3/1	1	not-connected	100	32	disabled
3/2	1	forwarding	100	32	disabled

Le processus de création de l'arbre spanning tree peut durer plusieurs dizaines de secondes. Ce temps est, en réalité, proportionnel au nombre de commutateurs.

Pendant cette phase, aucun commutateur ne traite de trame au cours des 15 premières secondes (valeur par défaut) ; le réseau s'arrête donc de fonctionner chaque fois qu'un commutateur est allumé ou éteint quelque part dans le réseau.

La phase d'apprentissage est encore plus longue lorsque tous les commutateurs s'initialisent en même temps (suite à une panne de courant, par exemple). En effet, les BPDU sont reçus par plusieurs ports dont l'un peut être élu racine, puis invalidé par la suite si un commutateur situé en aval a invalidé sa route. Le temps de stabilisation de l'arbre peut ainsi atteindre plusieurs minutes.

Dans certains cas, notamment sur les commutateurs fédérateurs, il peut être intéressant de diminuer ce temps, surtout si l'architecture réseau est conçue sans aucune boucle.

```
set spanntree fwwdelay 5
```

Inversement, si ce temps est trop court par rapport au délai de construction de l'arbre, des trames peuvent commencer à circuler et potentiellement être dupliquées dans le cas de routes multiples. Il vaut alors mieux augmenter le paramètre "forward delay" au-delà des 15 secondes par défaut.

La meilleure solution consiste à activer plus rapidement les ports qui ne sont pas concernés par le *spanning tree*, c'est-à-dire ceux sur lesquels sont connectés à une seule station.

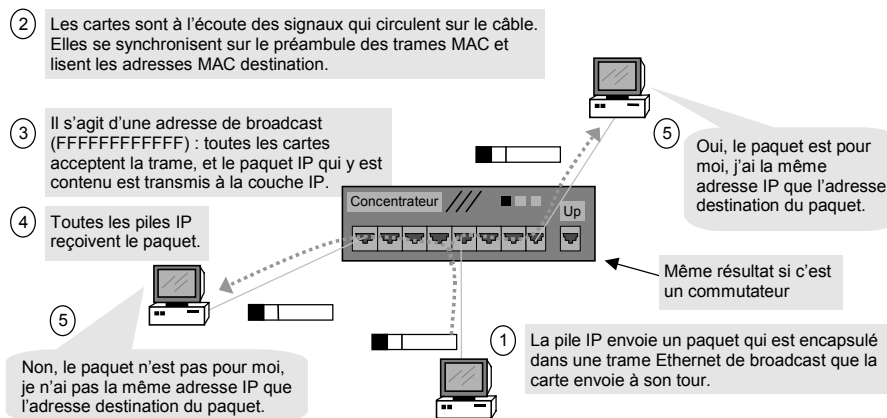
```
set spanntree portfast 1/2 enable
```

On peut remarquer que l'échange de BPDU et l'élection d'un commutateur désigné impliquent que chaque port du commutateur soit identifié par une adresse MAC (comme une carte réseau).

Comment une station envoie-t-elle un paquet IP à une autre ?

Une carte réseau ne se préoccupe que des adresses MAC pour envoyer et recevoir des données. En revanche, une application telle que Telnet ne connaît que l'adresse IP qui est purement logique : une pile IP recevant un paquet IP ne le prendra en compte que si l'adresse de destination du paquet correspond à l'adresse IP qui a été paramétrée dans le PC. Dans le cas contraire, il sera ignoré.

Figure 7-15.
*Échange
de paquets IP.*



Par ailleurs, l'adresse MAC de la station changera si la carte réseau est changée (en cas de panne, par exemple). De même, son adresse IP peut être modifiée à tout moment à l'aide des outils de configuration Windows (en cas de déménagement, par exemple).

L'exemple précédent montrait un paquet IP envoyé dans une trame de broadcast. Ce moyen d'opérer est pratique mais très consommateur de bande passante puisque la trame est propagée à travers tout le réseau. Sauf quand cela est nécessaire, un paquet IP est envoyé dans une trame unicast, c'est-à-dire directement au PC concerné. Mais comment connaître l'adresse MAC de la carte du PC destinataire alors que vous ne connaissez que l'adresse IP de sa pile IP ?

Cela est, par exemple, le cas lorsque vous lancez la commande suivante, qui permet de vous connecter à un serveur Unix :

```
Telnet 192.50.10.1
```

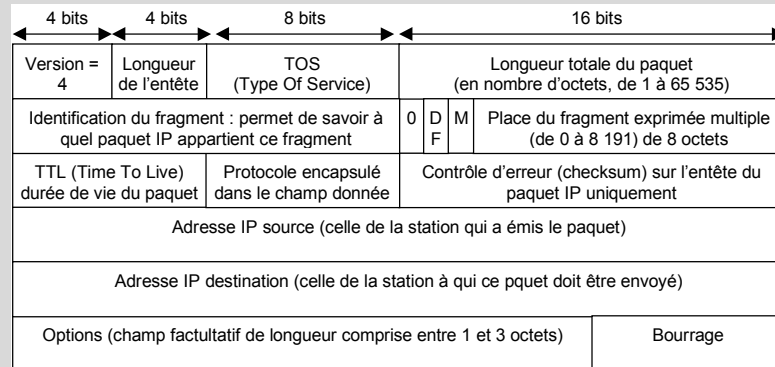
L'application Telnet va demander à la couche TCP d'ouvrir une connexion avec l'adresse IP indiquée, et va transmettre son paquet à la couche IP (avec l'adresse de destination indiquée). Cette dernière va encapsuler le paquet TCP dans un paquet IP, puis l'envoyer à la carte. Mais la carte ne sait pas quoi faire d'une adresse IP ; elle ne sait gérer que des adresses MAC : une trame Ethernet ne contient qu'une adresse MAC qui permet aux autres cartes de la prendre ou non en compte.

La solution repose sur un mécanisme qui réalise la correspondance entre l'adresse MAC du PC destinataire et son adresse IP.

On pourrait utiliser une table de correspondance statique Adresse MAC ↔ Adresse IP. Mais cela serait fastidieux, car il faudrait relever les adresses MAC des stations ainsi que les adresses IP, et paramétrer la table sur tous les PC. Cela est inimaginable étant donné le nombre important de PC et les nombreux changements d'adresses qui interviennent. On perdrait en plus l'avantage de dissocier l'adresse physique de l'adresse logique. En outre, un PC peut être configuré avec plusieurs adresses IP.

LE POINT SUR IP v4 (RFC 791)

IP (*Internet Protocol*) est un protocole de niveau 3 (couche réseau) qui découpe les réseaux locaux en réseaux logiques indépendamment de leur implémentation physique. Ce protocole permet donc d'envoyer des données à travers **un réseau virtuel** reposant sur des réseaux physiques de différente nature (Ethernet et PPP, par exemple). Pour ce faire, IP utilise un **adressage logique** différent de l'adressage physique (MAC, PPP ou autre).



Cette couche se contente de **router** (c'est-à-dire acheminer) le paquet à travers un réseau IP : les paquets peuvent être perdus (pas de garantie d'acheminement), contenir des erreurs (sauf sur l'en-tête, qui est contrôlé) ou arriver dans le désordre. IP **fragmente** les paquets dont la taille excède celle des trames (le MTU, *Maximum Transfer Unit*). Les fragments sont routés indépendamment les uns des autres comme autant de paquets, mais IP **assemble** dans le bon ordre les fragments d'un même paquet original. Si le bit "DF" est positionné à 1, la fragmentation est interdite. Le bit "M" positionné à 0 indique que ce paquet est le dernier fragment d'une série lorsque le bit "DF" est positionné à 0.

Le champ **TTL** est décrémenté de 1 chaque fois que le paquet passe par un routeur. Si la valeur atteint 0, le routeur détruit le paquet. Ce mécanisme évite aux paquets de rester trop longtemps sur le réseau, soit parce qu'ils tournent en boucle (suite à une erreur de routage), soit parce qu'ils traversent trop de routeurs. La valeur initiale du TTL est fixée par la station émettrice (de 32 à 128, en général).

D = délai d'acheminement court T = débit élevé R = Grande fiabilité					C = Option recopiée dans tous les fragments			Classe / Numéro		
								0 / 2 IP security Option		
								0 / 3 Routage lâche		
								0 / 7 Enregistrement des routes		
								0 / 9 Routage strict défini par la source		
								2 / 4 Horodatage des paquets		
TOS	Priorité	D	T	R	0	Options	C	Classe	Numéro	

Le champ **TOS** permet de décrire la qualité de service souhaitée. La signification de ce champ est abordée au chapitre 14.

La résolution d'adresse

La seule solution est donc une résolution d'adresse automatique, c'est-à-dire un mécanisme permettant de trouver l'adresse MAC en connaissant uniquement l'adresse IP. C'est le rôle du protocole **ARP** (*Address Resolution Protocol*) lié à la couche IP. Ce protocole gère une table de correspondance dynamique Adresse MAC ↔ Adresse IP, appelée **cache ARP**. Vous pouvez en visualiser le contenu à l'aide de la commande Windows suivante :

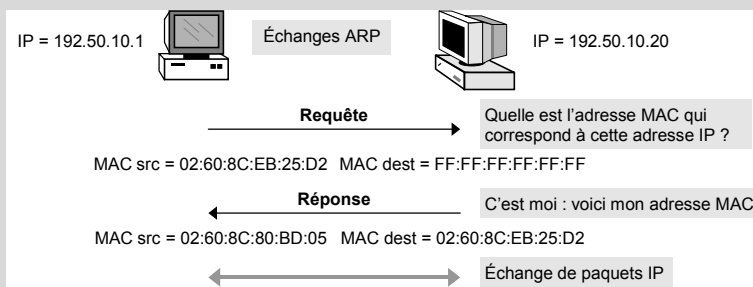
```
arp -a
00:40:0b:4b:25:d2    190.50.1.253
```

LE POINT SUR ARP (RFC 826)

Pour obtenir l'adresse MAC d'une station ne connaissant que son adresse IP, la pile TCP/IP émet une requête ARP (*Address Resolution Protocol*) dans une trame Ethernet de broadcast dont le champ « Type » contient la valeur **0x0806**.

Chaque pile IP recevant un tel paquet compare alors son adresse avec celle figurant dans le champ « Adresse protocole destination ».

S'il y a correspondance, la couche ARP envoie un paquet de réponse en remplissant le champ « Adresse physique destination » avec l'adresse MAC de sa carte. Dans le cas contraire, le paquet est ignoré.



Donc, seule la station dont l'adresse IP correspond à celle demandée par la requête envoie en réponse un paquet contenant sa propre adresse MAC.

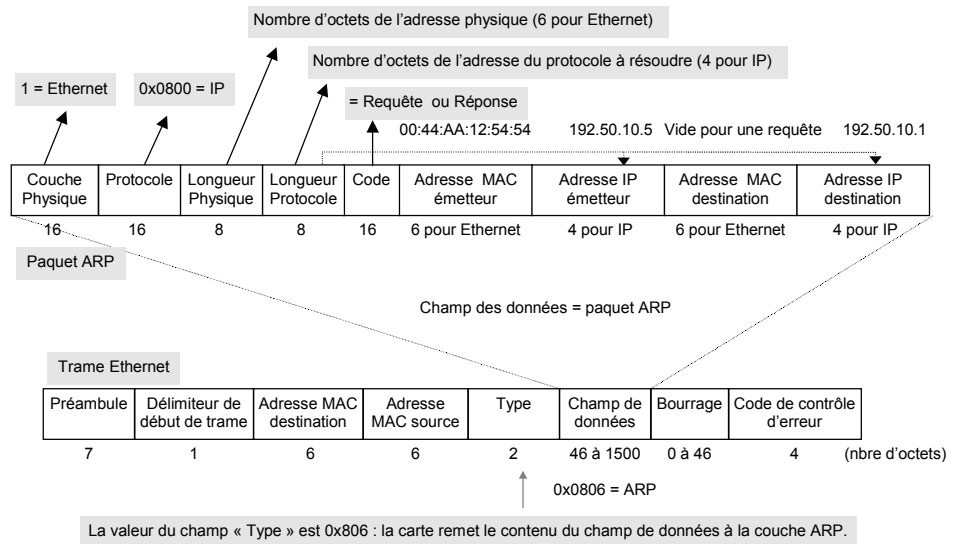
La résolution inverse, c'est-à-dire l'obtention de l'adresse IP à partir de l'adresse MAC, est réalisée par le protocole **RARP** (*Reverse ARP* – RFC 903).

Si vous n'avez pas communiqué récemment avec un autre PC, la table sera vide : les entrées sont, en effet, effacées au bout d'un certain temps. Sous Windows, une entrée ARP (adresse MAC / adresse IP) est supprimée au bout de deux minutes si le PC n'a pas dialogué avec la station cible (selon le mécanisme TTL, *Time To Live*). Dans tous les cas, l'entrée reste au maximum dix minutes en mémoire, puis elle est supprimée.

Si l'adresse IP recherchée n'est pas dans le cache, ARP va alors envoyer un paquet de requête encapsulé dans une trame Ethernet de broadcast. Cette dernière va donc être lue par toutes les cartes réseau.

Figure 7-16.

*Format
d'un paquet ARP
dans une trame
Ethernet.*



Seule la station configurée avec l'adresse IP recherchée va répondre en renvoyant son adresse MAC.

Une fois l'adresse résolue, le paquet IP peut être envoyé dans une trame MAC unicast dont l'adresse de destination est celle de la station cible.

Comment une application envoie-t-elle des données ?

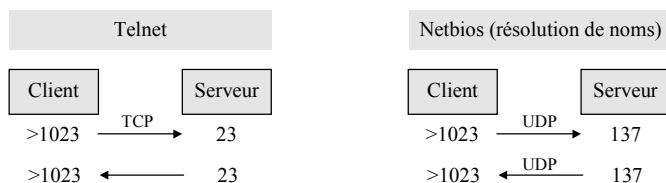
Une application utilise les services de la couche transport avec qui elle échange des données à travers une interface de programmation livrée avec la pile TCP/IP. Sous Unix, il s'agit des Sockets ; sous Windows de Winsock.

La couche transport est soit **TCP** (*Transport Control Protocol*) soit **UDP** (*User Datagram Protocol*), qui est une version allégée de TCP.

Le protocole TCP agit en **mode connecté**, ce qui implique que le client demande l'ouverture d'une connexion préalablement à tout échange. Par exemple, lorsque vous entrez la commande Windows "Telnet 192.50.10.1", le programme client Telnet demande à TCP d'ouvrir une connexion à un serveur Telnet qui est en attente, c'est-à-dire à l'écoute du **port** TCP 23.

Inversement, UDP agit en mode **non connecté**, ce qui permet à deux machines d'échanger des données à tout moment, sans entrer dans une phase de connexion. Par exemple, lorsque vous voulez vous connecter à un serveur de fichiers Windows NT, votre PC émet une demande de résolution de nom à destination d'un serveur WINS qui est à l'écoute sur le **port** UDP 137.

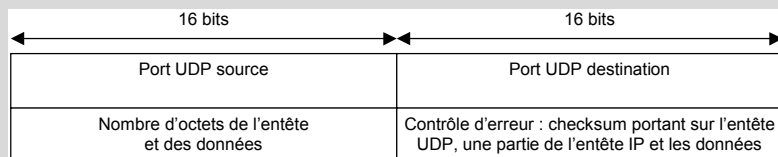
Figure 7-17.
Utilisation
des ports TCP et UDP.



La pile TCP/UDP du client choisit généralement un port source supérieur à 1023, et incrémente cette valeur à chaque nouvelle session ouverte simultanément à d'autres déjà actives.

LE POINT SUR UDP (RFC 768)

UDP (*User Datagram Protocol*) permet simplement à une application d'avoir accès au réseau IP. Ce protocole n'offre aucune garantie d'acheminement, aucun mécanisme de reprise sur erreur, ni de contrôle de flux, et ne vérifie pas la duplication des paquets. Tous ces contrôles doivent être opérés par les autres couches réseau. En revanche, les paquets remis à l'application le sont sans erreur.



Les champs **port source** et **port destination** servent à identifier une application (par exemple, 23 pour Telnet). Ces valeurs sont réservées et enregistrées par l'IANA (*well known port* – RFC 1700).

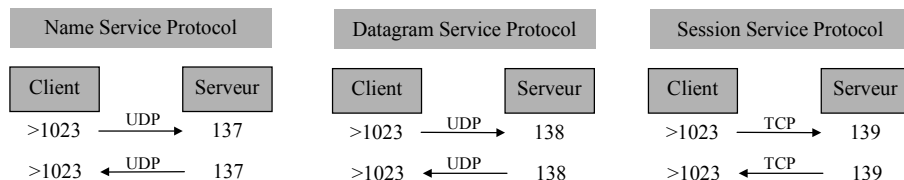
Les clients et les serveurs manipulent des noms (une machine Unix, un serveur de fichiers NT) et s'échangent des données à travers des ports qui leur sont réservés par l'IANA (voir chapitre 3). Sur l'Internet, le service DNS permet de convertir les noms en adresses IP (voir chapitre 17), tandis que, dans le monde Microsoft, on utilise encore le service WINS pour convertir des noms **Netbios** en adresses IP.

On peut considérer que ce protocole est situé au niveau de la couche 5 (couche session) : il permet, en effet, d'établir et de gérer des sessions entre applications. Dans le monde Internet, les applications comme Telnet, votre navigateur web, FTP, etc.) gèrent elles-mêmes tous les mécanismes situés au-dessus de la couche transport, c'est-à-dire TCP et UDP.

À l'origine, Netbios circulait nativement dans des trames Ethernet, mais de nos jours, il est encapsulé dans IP (RFC 1001 et 1002).

Par exemple, le partage de fichiers et la messagerie Exchange utilisent le protocole Netbios sur le port TCP 139. Par ailleurs, les serveurs WINS s'échangent des données sur le port TCP 42.

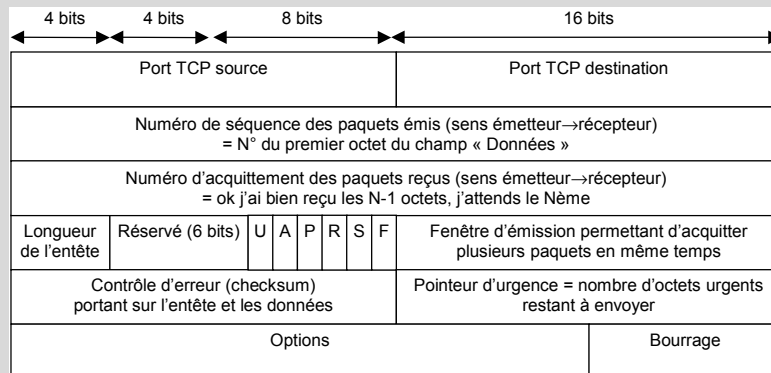
Figure 7-18.
Netbios sur IP.



LE POINT SUR TCP (RFC 793)

TCP (*Transport Control Protocol*) est un protocole de niveau 4 (couche transport) qui permet à deux applications (un client et un serveur) d'échanger des données en leur masquant les mécanismes réseau. Les paquets TCP sont transportés dans des paquets IP de type 6 (champ protocole = 6).

TCP offre un service de **bout en bout** (entre deux entités, quelle que soit leur localisation) en **mode connecté** (un client doit se connecter à une application serveur). Il utilise pour cela un adressage applicatif basé sur des **ports TCP**. Chaque application est identifiée par un numéro de port réservé (*well known port*). Généralement, le client choisit un numéro de port aléatoire supérieur à 1023 comme port source. Le serveur lui répond sur ce port.



Les bits de contrôle U, A, P, R, S et F ont la signification suivante :

U = Urgent. Indique que les données doivent être remises sans délai à l'application.

A = Ack. Acquitement d'une demande de connexion ou de fermeture.

P = Push. Indique à la couche TCP d'envoyer et de remettre les données sans attendre le remplissage des tampons d'émission et de réception.

R = Reset. Ferme la connexion TCP suite à un problème.

S = Synchrhonize. Le numéro de séquence est réinitialisé à une valeur aléatoire.

F = Fin. Demande de déconnexion.

Des options peuvent être négociées entre entités TCP (champ " Option "), par exemple la taille maximale des segments transportés.

La couche TCP assure le **contrôle d'erreur** et le **séquencement** des paquets (les paquets sont remis dans le même ordre que lors de leur émission). La taille de la **fenêtre d'émission** indique le nombre de paquets pouvant être acquittés en même temps. Elle permet également de demander la retransmission à partir du premier paquet en erreur (manquant ou erroné).

La couche TCP mesure le temps écoulé entre l'émission d'un paquet et la réception de l'accusé de réception correspondant, et calcule ainsi une moyenne glissante du temps de réponse (*Round Trip Time*). Elle utilise l'algorithme de Karn pour déduire la valeur de ses temporisateurs. Par exemple, plus le temps de réponse est long, plus TCP attendra longtemps l'accusé de réception avant de retransmettre. De même, TCP estime le nombre de paquets perdus : plus celui-ci augmente, plus la fenêtre d'émission est réduite. Ces mécanismes permettent à TCP de **contrôler le flux** de données en fonction de l'état du réseau (perte de paquets et débits) et donc d'éviter une surcharge du réseau par un nombre croissant de retransmissions devenues inutiles.

Mettre en place sa première interconnexion de réseaux

Comme leur nom l'indique, les réseaux locaux sont géographiquement restreints à un immeuble, voire à un campus. Votre société se développant, de nouveaux sites sont créés et les réseaux locaux se multiplient.

L'enjeu est désormais de connecter ces réseaux entre eux de sorte que tous les utilisateurs accèdent aux mêmes applications quelle que soit leur localisation.

C'est le rôle des réseaux WAN (*Wide Area Network*), c'est-à-dire des réseaux étendus. On parle également d'interconnexion de réseaux, de réseaux longue distance ou de réseaux intersite.

Dans ce chapitre, vous apprendrez :

- à choisir entre plusieurs solutions techniques et économiques ;
- à interconnecter deux sites ;
- à mettre en place une ligne spécialisée et une ligne de secours RNIS ;
- à configurer un routeur.

Le contexte

Notre réseau local est opérationnel, et les utilisateurs sont satisfaits. En plus du site parisien, il faut maintenant offrir le même type de service pour un nouveau site situé à Orléans. Qu'à cela ne tienne, il suffit d'appliquer de nouveau les recettes qui ont déjà fait le succès de notre premier réseau.

Mais les utilisateurs de chaque site doivent communiquer entre eux : messagerie, transferts de fichiers et connexions aux serveurs web sont demandés. Il faut donc interconnecter les réseaux locaux de ces deux sites. Le problème est qu'ils sont distants de plus de 100 km.

On pourrait utiliser les mêmes équipements Ethernet, des commutateurs par exemple. Mais la norme impose une longueur maximale aux liaisons Ethernet, au mieux quelques kilomètres en fibre optique. Ces contraintes proviennent de l'affaiblissement du signal d'une part, et du délai de propagation des trames, d'autre part. En effet, plus les distances sont grandes, plus le signal est affaibli et plus le délai de propagation des trames est élevé. Si ce dernier était plus élevé que celui imposé par la norme, les stations devraient attendre plus longtemps avant de pouvoir transmettre une trame, ce qui diminuerait considérablement le débit du réseau (moins de trames circuleraient en un laps de temps donné puisqu'il faudrait attendre plus longtemps avant de transmettre). À l'avenir, cette contrainte disparaîtra mais, pour le moment, nous devons encore en tenir compte.

Il faut donc employer d'autres techniques plus adaptées à ces contraintes et définir une architecture des réseaux étendus, appelés **WAN** (*Wide Area Network*) par opposition aux réseaux locaux, LAN (*Local Area Network*).

Les choix de base

Quel support de transmission ?

Le RTC (réseau téléphonique commuté) que nous avons utilisé pour nous connecter à l'Internet est un réseau étendu.

Quelques supports de transmission utilisés pour les réseaux étendus		Débit
RTC	Réseau téléphonique analogique utilisé pour transporter des données.	De 19,2 à 56,6 Kbit/s
RNIS (réseau numérique à intégration de services)	Réseau téléphonique numérique utilisé pour transporter des données. Très utilisé en interconnexion de LAN.	De 64 à 128 Kbit/s (plus rare : N x 64 Kbit/s)
LS (ligne spécialisée)	Liaison numérique en point à point entre deux sites. Très utilisée en interconnexion de LAN.	De 64 Kbit/s à 2 Mbit/s ou 34 Mbit/s

Quelques supports de transmission utilisés pour les réseaux étendus		Débit
xDSL (<i>x Digital Subscriber Line</i>)	Liaisons numériques en point à point entre deux sites. De plus en plus utilisées pour les accès à l'Internet.	De 64 Kbit/s à 6 Mbit/s
SDH (<i>Synchronous Data Hierarchy</i>)	Liaisons en fibre optique à haut débit utilisées par les opérateurs.	De 51 Mbit/s à plusieurs Gbit/s
ATM (<i>Asynchronous Transfert Mode</i>)	Liaisons en fibre optique à haut débit (emprunte également des supports SDH). ATM est également utilisé pour les LAN.	155, 622 Mbit/s et plus
Frame Relay	Liaisons numériques à commutation de trames. Très utilisées en interconnexion voix/données.	De 64 Kbit/s à 34 Mbit/s

Ces supports de transmission nécessitent des modems adaptés dénommés CSU (*Channel Service Unit*) ou DCE (*Data Circuit-Terminating Equipment*), par opposition aux équipements DSU (*Data Service Unit*) ou DTE (*Data Terminal Equipment*) qui s'y connectent. Par exemple, le modem est un DCE, et le PC un DTE. Nous nous situons ici au niveau physique.

Quel protocole de niveau 2 ?

Les concepteurs auraient pu utiliser les mêmes trames Ethernet, et ainsi simplifier le problème, d'autant que la possibilité d'adressage est immense (2^{48} adresses). Cela aurait pu être le cas, mais il aurait fallu élaborer un mécanisme spécifique pour ne pas propager dans tout le réseau les trames de broadcast et multicast. Cela aurait entraîné d'autres complications inextricables puisqu'il faut quand même les propager dans un certain périmètre (résolution d'adresses, etc.). De plus, une trame Ethernet comprend 18 octets, ce qui était considéré comme un overhead important lorsque le débit des réseaux étendus était limité (ce qui est d'ailleurs toujours le cas avec notre modem RTC). En définitive, les protocoles LAN ne sont pas adaptés aux réseaux étendus. Ce constat sera de moins en moins valable dans le futur.

Quelques protocoles de niveau 2 utilisés pour les réseaux étendus	
PPP (Point to Point Protocol)	Utilisé sur les supports RTC, RNIS, LS et ADSL.
Frame-Relay	Protocole point à point et multipoint voix et données. Utilisé sur des supports LS.
ATM (Asynchronous Transfert Mode)	Comme pour Ethernet, la norme définit les couches physique et liaison. ATM véhicule voix et données.

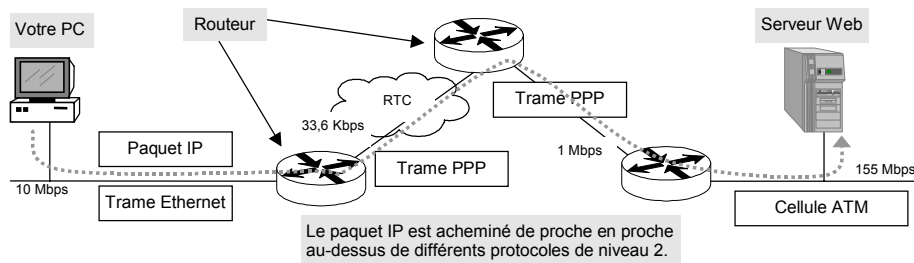
Les deux tableaux précédents font apparaître ATM comme étant le protocole universel : il fonctionne sur les LAN et les WAN et supporte la voix et les données. Mais, bien que très utilisé par les opérateurs sur leur réseau WAN, il est très peu utilisé en LAN à cause de son coût.

Quel équipement réseau ?

L'utilisation de multiples protocoles de niveau 2 pour transporter les paquets IP pose un nouveau problème en comparaison des architectures de réseaux locaux qui n'utilisaient qu'Ethernet. Celui-ci avait été entr'aperçu lorsque l'encapsulation des paquets au-dessus de plusieurs réseaux avait été étudiée au chapitre 7.

Figure 8-1.

L'acheminement d'un paquet IP.



Comment, en effet, assurer la continuité d'adressage et de commutation au-dessus de protocoles aussi différents ? Réponse : le paquet IP est le seul lien commun. Il faut donc disposer d'équipements spécifiques qui permettent de :

- gérer les différents supports de transmission LAN et WAN ;
- traiter les paquets IP, c'est-à-dire utiliser les protocoles de niveau 3.

L'équipement qui répond à ces besoins est le **routeur**, c'est-à-dire un commutateur de niveau 3 (par opposition aux commutateurs de niveau 2, tel que trouvés sur Ethernet).

Quel opérateur ?

De plus, notre interconnexion de réseaux doit passer dans des zones du domaine public dont nous n'avons pas la maîtrise. Enfin, même en admettant que nous obtenions toutes les autorisations administratives nécessaires, la pose de câbles entre les deux sites reviendrait très cher.

La seule solution est de faire appel aux services d'un opérateur tel que France Télécom, Cegetel, Colt, etc. Dans ce domaine, le marché offre un nombre impressionnant de solutions combinant techniques et niveaux de service.

Niveau de prestation	Description technique	Service fourni
1. Support de transmission (couche physique)	LS, xDSL, ATM (connexions point à point) et RNIS (multipoint)	Supervision de la ligne (option : garantie de temps de réparation)
2. Réseau fédérateur (couches physique et liaison)	Accès via LS et RNIS au backbone Frame-Relay, ATM, etc., de l'opérateur	Réseau fourni et exploité par l'opérateur + support client avec engagements de résultats
3. Interconnexion de réseaux locaux (couches 1, 2 et 3)	Support de transmission + réseau fédérateur + routeur	Réseau étendu de bout en bout fourni et exploité par l'opérateur + support client, avec engagement de résultat

Les coûts associés à ces services sont de différentes natures :

- frais uniques de mise en service ;
- frais mensuels fixes en fonction du débit des lignes et de la qualité du service ;
- et, de moins en moins, frais mensuels variables en fonction de la consommation.

De quoi avons-nous besoin ?

Pour notre première interconnexion, nous allons restreindre notre choix au plus simple et au moins cher.

D'une liaison entre les deux sites

Les sites de Paris et d'Orléans étant distants de 112 km à vol d'oiseau, il faut obligatoirement passer par un opérateur. Mais quel support de transmission utiliser et à quel débit, et quel type de service demander ?

À quel débit ?

Le dimensionnement des liaisons est un exercice délicat et important, car il va influencer sur les temps de réponse du réseau et donc sur la satisfaction des utilisateurs. À l'inverse des LAN, les débits des réseaux étendus sont limités à cause des coûts qu'ils entraînent.

Le débit dépend de trois facteurs : le type de trafic, le volume de données généré et les temps de réponse requis.

Application	Type de trafic	Volume généré	Temps de réponse requis
Sauvegarde	Transfert de gros et très gros fichiers	Élevé	Faible à moyen
Partage de fichiers avec des serveurs Windows NT	Transfert de petits et gros fichiers	Élevé	Moyen
Messagerie	Transfert de petits et gros fichiers	Moyen à élevé	Moyen à élevé
Base de données client-serveur	Transactionnel (transfert de petits et moyens fichiers)	Faible à élevé	Moyen
Connexion aux serveurs web	Transactionnel (transfert de petits et moyens fichiers)	Faible	Moyen
Connexion Telnet sur un serveur Unix	Conversationnel (écho distant)	Faible	Élevé (< 300 ms)

Les transferts de fichiers perturbent fortement les flux conversationnels (Telnet) et, dans une moindre mesure, les petits flux transactionnels (serveurs web). En conséquence, les sauvegardes et les transferts de gros fichiers doivent, de préférence, être effectués la nuit. Cela présente le double avantage de ne pas gêner le travail des utilisateurs de jour et de répartir l'utilisation du réseau sur 24 heures.

Dans un cas simple comme le nôtre (messagerie, transfert de fichiers et connexion web), un débit de 64 à 128 Kbit/s devrait être suffisant. Compte tenu du nombre d'utilisateurs (plusieurs centaines sur chaque site), un débit de 128 Kbit/s est plus sécurisant. Nous limiterons les risques en souscrivant un contrat d'une durée minimale d'un an, ce qui nous permettra de changer de débit facilement.

Lorsque notre réseau intersite aura une plus grande ampleur, nous devons nous livrer à un calcul plus précis, comme nous le verrons au chapitre 9.

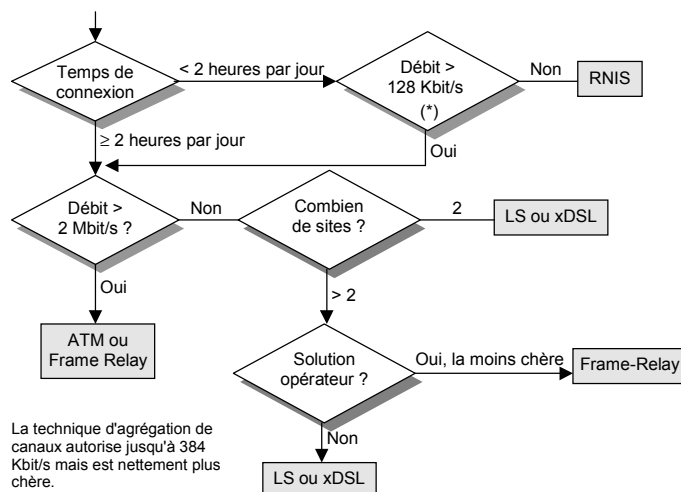
Avec quel support de transmission ?

La manière de calculer les frais mensuels varie d'un support de transmission à l'autre et, parfois, d'un opérateur à l'autre.

Support de transmission	Mode de facturation	Frais mensuels
RNIS	Durée de la communication et distance	Variables
LS / xDSL	Débit et distance	Fixes
Frame Relay	Débit garanti, débit du port et type d'équipement	Fixes
ATM	Débit et qualité de service	Fixes

Le premier critère à prendre en compte est donc le temps d'utilisation de la liaison, puis son débit.

Figure 8-2.
*Choix d'un support
pour une interconnexion
de réseaux locaux.*



NOTE Le seuil de deux heures par jour a été calculé pour une distance de 100 km.

Dans notre cas, la liaison devrait être utilisée toute la journée, quasiment en continu compte tenu du nombre d'utilisateurs (quelques centaines) et des applications utilisées (flux client-serveur, connexions à des machines Unix, etc.). La LS point à point est donc appropriée.

Avec quel service opérateur ?

En ce qui concerne le type de service, la question est de savoir si vous voulez « faire » ou « faire faire ». En d'autres termes, voulez-vous réaliser vous-même l'interconnexion (solution privée), ou préférez-vous confier le projet à un spécialiste (solution opérateur).

Le premier critère à prendre en compte est celui de l'implication en termes de ressources humaines.

Tâches \ Qui fait ?	Solution privée	Service opérateur niveau 3
Étude technique (choix des supports, dimensionnement, etc.)	Vous	Vous, assisté de l'opérateur
Fourniture de la liaison	Opérateur	Opérateur
Achat des routeurs	Vous, auprès d'un distributeur réseau	Opérateur
Installation et configuration des routeurs	Vous ou le distributeur réseau	Opérateur
Exploitation du routeur (supervision, modification des paramètres, intervention lors des pannes, etc.)	Vous ou le distributeur, avec un contrat d'assistance	Opérateur
Résolution des pannes matérielles	Distributeur, avec un contrat de maintenance	Opérateur

La solution opérateur nécessite moins de ressources internes (vous êtes notamment moins sollicité), surtout avec le service de niveau 3. La solution privée revient à choisir le service de niveau 1 et implique donc de gérer un projet avec au moins deux interlocuteurs (le distributeur et l'opérateur).

Le second critère à prendre en compte est d'ordre financier.

Tâches \ Quel coût ?	Solution privée	Service opérateur niveau 3
Étude technique (choix des supports, dimensionnement, etc.)	Dépend de l'importance du projet	Coût identique
Mise en œuvre du support de transmission	Coût standard (mise en service + frais mensuels)	Coût moins élevé, car la LS d'accès est locale
Achat des routeurs	Coût d'investissement	Frais mensuels d'exploitation
Installation et configuration des routeurs	Coût d'investissement	Frais de mise en service
Exploitation	Contrat d'assistance sur site	Coût plus élevé, mais meilleure qualité de service
Résolution des pannes matérielles	Contrat de maintenance annuel : 8 à 12 % du coût d'achat	Moins cher

La solution opérateur devrait être la moins chère à long terme (deux à cinq ans) et apporter le moins de soucis. Cela n'est cependant pas toujours le cas, surtout pour de petits réseaux, qui plus est, limités à la France. Les économies les plus importantes sont, en effet, réalisées à l'international.

Si vous disposez d'un gros réseau, vous pouvez opter pour la solution intermédiaire qui consiste à retenir le service opérateur de niveau 2. Cela permet de conserver la maîtrise du routage IP et de modifier plus facilement la configuration de vos routeurs. Rien ne vous empêchera par la suite d'opter pour le service de niveau 3 (tous les opérateurs proposent des routeurs Cisco qui peuvent donc être repris).

Dans notre cas (deux sites à interconnecter), la meilleure solution est une LS avec une garantie de temps de rétablissement. Mais il serait intéressant de comparer les coûts avec un service de niveau 3 (LS plus routeur). Les résultats de cette comparaison peuvent varier en fonction de nombreux critères (le moment où elle est réalisée – car les prix évoluent vite –, les distances entre sites, le débit, etc.).

Admettons cependant que la solution privée soit la moins chère ou alors regardons faire l'opérateur qui fait ce que nous aurions dû faire.

De routeurs

Le routeur est le seul équipement permettant d'interconnecter deux sites sur de longues distances. Le réseau Internet n'est d'ailleurs constitué que de routeurs utilisant des liaisons spécialisées, ATM et Frame Relay.

Le plus simple est d'opter pour une configuration fixe comprenant une interface Ethernet et une interface WAN et ne supportant que le protocole IP. Le coût de cette configuration de base oscille entre 7 000 et 9 000 F HT. Il nous faut deux routeurs, un par site.

Quelques extensions matérielles et logicielles sont proposées.

Fonctionnalité	Description	Intérêt	Coût HT
Multiprotocole	Support des protocoles IP + IPX + Decnet + SNA, etc.	Si existant à supporter	± 10 000 F (nécessite davantage de mémoire)
Fonction pont (bridge)	Commutation de niveau 2	Pour les anciens protocoles non routables	± 1 000 F (ou 0 F car souvent livrée en standard)
Interface RNIS	Connexion au support de transmission RNIS	Peut être utilisée pour le secours ou le débordement	± 3 000

Dans notre cas, l'option RNIS est intéressante, car elle offre une solution de secours en cas de panne de la LS. Cela permet ainsi d'assurer la continuité de service avec la même qualité, le débit offert étant de 64 et 128 Kbit/s, soit exactement celui de notre LS.

L'autre intérêt du RNIS concerne le débordement : lorsque la LS est chargée à 80 ou 100 %, le routeur peut activer l'interface RNIS, offrant ainsi un débit supplémentaire de 64 Kbit/s, voire 128 Kbit/s. Cette fonction pourra être utilisée si le débit nécessaire a été sous-évalué.

QU'EST-CE QU'UN PONT ?

Le **pont** (bridge) est l'ancêtre du commutateur de niveau 2. Il a été le premier équipement utilisé pour **interconnecter** des segments Ethernet partagés, soit localement soit *via* des réseaux étendus. Compte tenu de la technologie et des prix d'alors, il était équipé de deux à quatre interfaces LAN et WAN, et offrait une faible puissance comparé aux commutateurs d'aujourd'hui.

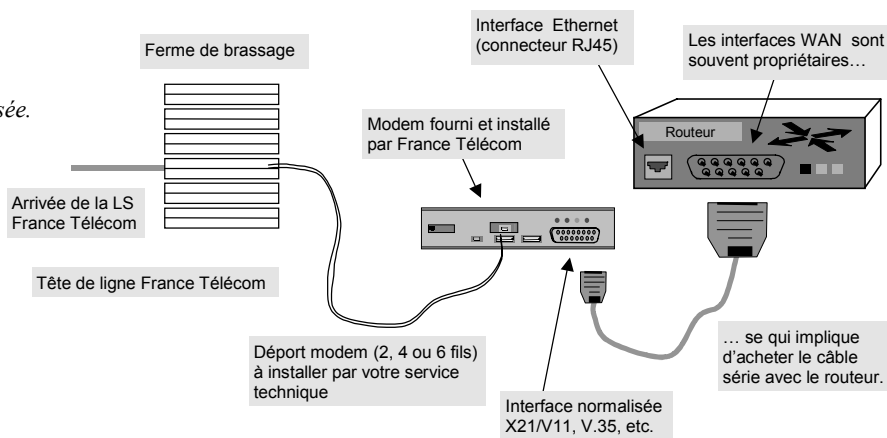
Les routeurs continuent de supporter cette fonctionnalité afin d'assurer la compatibilité avec l'existant.

De câbles

Encore des câbles ! Il faut bien sûr un cordon de brassage pour connecter le routeur à un concentrateur ou à un commutateur Ethernet. L'interface Ethernet du routeur étant strictement équivalente à une carte réseau d'un PC, un cordon identique (droit RJ45/RJ45) sera utilisé.

Il faut également des câbles spécifiques côté réseau de l'opérateur, en l'occurrence côté ligne spécialisée. Pour notre connexion à l'Internet, nous avons utilisé un câble série entre le PC et le modem, câble fourni avec le modem. Le principe est identique pour les routeurs : lorsque l'opérateur (France Télécom, par exemple) met en service la LS, il installe dans vos locaux un modem auquel vous connectez le routeur *via* un câble série adapté.

Figure 8-3.
*Connexion
du routeur
à la ligne spécialisée.*



La ligne spécialisée arrive dans un local technique réservé à France Télécom. La longueur des câbles série étant limitée, il faut déporter le modem au plus près du routeur qui se trouve dans un LTE ou une salle informatique. Cette partie de l'installation est privée (elle se déroule dans les locaux de votre société) et doit donc être réalisée par les services techniques de l'immeuble (souvent, les services généraux ou les téléphonistes). France Télécom se contente d'amener la LS dans son local, d'attendre l'installation du déport modem est terminée, puis d'installer le modem et teste la LS avant de la mettre en service.

Les câbles série sont de différentes natures. Pour la connexion Internet, nous avons utilisé un câble V.24 (connexion série RS-232). Pour notre LS, l'opérateur France Télécom nous propose V.35 ou X21/V11.

Spécifications des interfaces					
Appellation usuelle	Mécanique	Électrique	Fonctionnelle	Longueur du câble	Débit en Kbit/s
V.24 (RS 232)	ISO 2110 25 broches	V.28	V.24	12 m	De 2,4 à 19,2
V.35 (RS 232)	ISO 2693 34 broches	V.11/V.10	V.24	15 m 10 m	48, 56, 64 128, 256
V.36	ISO 4902 37 broches	V.11/V.10	V.24	15 m 10 m	48, 56, 64 128, 256
X21/V.11	ISO 4903 15 broches	V.11	X.24	100 m 50 m	De 64 à 1 024 1 920
G703	ETSI 300.166	G703	G703	300 m	2 048
G703/G704	ETSI 300.167 9 broches	G703	G704	300 m	De 256 à 1 984



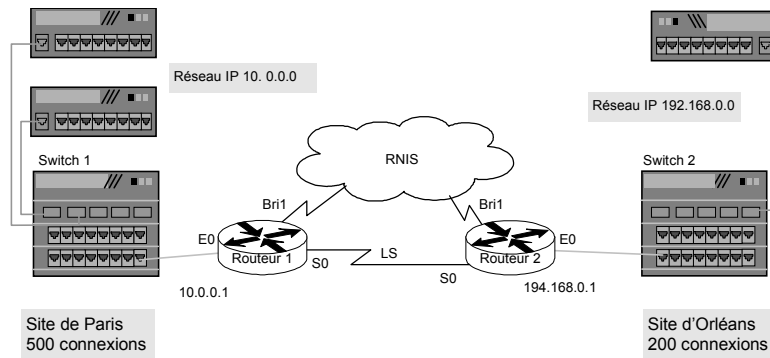
Malgré la normalisation, des variantes peuvent cependant exister. Les plus connues sont celles relatives à la norme V.35 qui accepte une version américaine et une version française en fonction du format de la prise. Les différences résident d'une part dans le diamètre des broches (plus grand sur le modèle américain) et d'autre part dans le type de fixation (par vis dans le premier cas, par clip dans le second).

Comment faire fonctionner tout cela ?

Définir l'architecture

Les routeurs vont s'insérer dans un réseau existant. Il convient donc de définir l'architecture et de réfléchir au paramétrage des routeurs.

Figure 8-4.
*Principe
de l'architecture réseau.*



Suivant les indications de notre plan d'adressage (voir chapitre 7), nous avons choisi pour notre site parisien une classe A « subnettée » sur 22 bits. En revanche, le site d'Orléans correspond à une nouvelle entité rachetée par la société, et une classe C avait été choisie par notre prédécesseur. Nous n'avons pas le temps de changer cette adresse et la conservons en l'état.

La question encore en suspend est de savoir si les interfaces WAN doivent également disposer d'une adresse IP et, dans l'affirmative, de quelle classe d'adresse.

Connecter un PC au routeur

Pour une configuration fixe, aucune préparation matérielle spécifique n'est nécessaire (pas de carte à installer, rien à démonter).

En revanche, le routeur est livré avec un logiciel non configuré (comme les PC). Vous pouvez le configurer dans votre bureau sans qu'il soit connecté au réseau, ou l'installer et le configurer sur place.

Les premiers paramétrages d'un routeur nécessitent de s'y connecter directement pour saisir des commandes. Un routeur est à l'image d'un PC : il contient un système d'exploitation (propriétaire) et une interface utilisateur (généralement en mode texte sans fenêtre ni souris).

Vous avez donc besoin d'un câble série (RS-232 avec un connecteur V.24) pour raccorder le port console du routeur au port série d'un PC.

Le type de port série des routeurs est variable. Vous pouvez trouver un connecteur DB9 comme sur votre PC ou, plus rarement, un connecteur DB25 (l'équivalent d'un DB9 avec 25 broches) ou encore, de plus en plus souvent, une prise RJ45 femelle. Il faut donc trouver le bon câble (croisé !) qui dispose d'un connecteur DB9 femelle côté PC (DTE) et d'un connecteur adéquat mâle côté routeur. De plus en plus souvent, ce cordon est livré avec le routeur.

Côté PC, vous devez utiliser le logiciel Hyperterminal en cliquant sur "Programme→Accessoires→Hyperterminal". Ce logiciel est un émulateur VT (*Virtual Terminal*) qui va vous permettre de dialoguer en mode texte avec le routeur.

Saisissez alors le nom du profil que vous pourrez réutiliser par la suite, par exemple “routeur”, puis cliquez sur “OK”. Choisissez ensuite le port série auquel vous avez connecté le câble reliant le PC au routeur : généralement COM1 ou COM2 (écran ci-contre).

Puis, saisissez les paramètres de la liaison série. Les valeurs dépendent du routeur employé ; elles sont généralement indiquées dans la documentation. Nos routeurs étant des routeurs Cisco, les valeurs sont les suivantes (deuxième écran ci-contre) :

Cliquez sur “OK”, puis appuyez plusieurs fois sur la touche “Entrée” ou, si cela n’est déjà fait, allumez le routeur.

Vous devez vous retrouver avec un écran de bienvenue et, un prompt vous invitant à saisir des commandes, ou bien vous obtenez un menu dans lequel vous vous déplacez à l’aide des touches fléchées du clavier.

Vous êtes maintenant prêt à configurer le routeur.

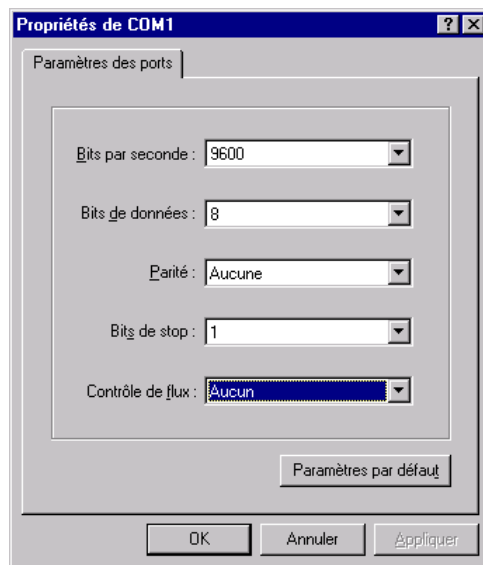
Configurer le routeur

Selon les constructeurs, la connexion réseau d’un routeur — l’équivalent d’une carte réseau d’un PC — s’appelle « port » ou « interface ». On parlera alors indifféremment de ports LAN ou WAN, d’interface série, etc.

Affecter les adresses IP

La première chose à faire est d’affecter les adresses IP à chaque interface. Côté LAN, on se réfère au plan d’adressage. La commande suivante réalise l’opération et active le logiciel IP dans le routeur.

```
int e 0
ip address 10.0.0.1 255.255.252.0
```



Abrégé de “Interface ethernet 0”

Par contre, faut-il en affecter une aux interfaces séries ? Si tel est le cas, il faut dédier un réseau IP pour seulement deux adresses. Même pour une classe C, 252 adresses sont gaspillées. Ceci étant, nous pouvons utiliser autant d'adresses que souhaité puisque nous avons pris le parti d'un plan d'adressage privé. Cependant, il est conseillé de limiter le nombre de réseaux afin d'en simplifier la configuration.

Deux solutions sont conseillées : aucune adresse IP (pour les petits réseaux), ou un réseau IP dédié aux interfaces WAN en le « subnettant » (pour les petits et grands réseaux).

Aucune adresse IP	Adresse IP « subnettée »
Point à point uniquement avec HDLC, PPP et Frame Relay en point à point.	Point à point et multipoint (Frame Relay, X.25, SMDS)
L'interface ne répond jamais au ping (pour savoir si elle est active) ; SNMP fonctionne.	L'interface répond au ping : cela permet de savoir rapidement si elle est active.
Autres fonctions propres au constructeur non supportées.	---

La commande suivante permet d'activer IP sur l'interface série sans lui affecter une adresse. Si des paquets sont générés par cette interface, l'adresse IP source sera celle de l'interface Ethernet.

```
int s 0  
ip unnumbered e 0
```

Abrégé de "Interface serial 0"

Par défaut, le routeur utilise un protocole de niveau 2, en général HDLC ou PPP. La commande suivante permet de forcer l'utilisation de PPP, qui est plus approprié aux liaisons point à point et à IP.

```
int s 0  
encapsulation ppp
```

Activer le routage

Le routage des paquets IP s'appuie sur la partie réseau des adresses de destination des paquets. Il faut donc indiquer au routeur parisien que les paquets à destination d'Orléans doivent être envoyés sur l'interface série.

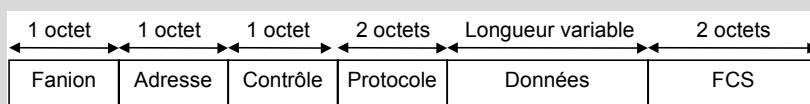
Dans notre cas, le plus simple est de configurer une route statique. Étant donné que notre interface série ne possède pas d'adresse IP, on indique explicitement son nom.

```
# Sur le routeur de Paris  
ip route 192.168.0.0 255.255.255.0 s0
```


LE POINT SUR PPP (RFC 1661 ET 1662)

PPP (*Point-to-Point Protocol*) est un protocole utilisé sur des liaisons point à point **synchrones** (telles qu'une LS) et **asynchrones** (avec un modem RTC, par exemple). Il est utilisé pour des débits variant entre 19,2 Kbit/s et 2 Mbit/s. Il est composé de trois éléments :

- des trames pour transporter les protocoles de niveau 3 ;
- de LCP (*Link Control Protocol*) pour établir, configurer et tester la liaison ;
- de NCP (*Network Control Protocols*) pour établir et configurer différents protocoles de niveau 3 (tels que IP, IPX ou Decnet) qui peuvent être multiplexés sur une seule liaison.



Le champ "Fanion" est un délimiteur de trame de valeur binaire "01111110". Le champ "Protocole" indique le type de protocole de niveau 3 présent dans le champ de données (0021 pour IP, c021 pour LCP, c023 pour LQM, c223 pour CHAP, etc.). Le champ "FCS" (*Frame Check Sequence*) est un code de détection d'erreur.

Dans sa version asynchrone, chaque octet de la trame est transmis avec un bit *start* et un bit *stop*, mais sans bit de parité, puisque le champ FCS est utilisé pour détecter les erreurs.

Le protocole **LCP** permet d'établir, de configurer, de surveiller et de terminer les liaisons point à point. Il permet, tout d'abord, de négocier des options, telles que :

- le MTU (*Maximum Transmission Unit*) dont la valeur par défaut est celui d'Ethernet (1500 octets) ;
- le choix du protocole d'authentification (CHAP, PAP ou aucun) ;
- le choix du protocole de contrôle de qualité (LQR ou aucun) ;
- la réduction du champ "Protocole" à un octet ;
- la suppression des champs "Adresse" et "Contrôle" lorsqu'ils ne sont pas utilisés.

LCP peut activer un protocole d'**authentification**, tel que **CHAP** (*Challenge Handshake Authentication Protocol*). Un mot de passe chiffré est échangé entre les deux nœuds (routeur ou PC).

LCP peut activer la procédure **LQM** (*Link Quality Monitor* — RFC 1989) qui coupe la liaison lorsque la qualité de service calculée tombe en dessous d'un seuil prédéfini. La qualité en émission est calculée en comparant le nombre total de paquets et d'octets transmis avec ceux reçus par la station distante. De même, la qualité en réception est calculée en comparant le nombre total de paquets et d'octets reçus avec ceux émis par la station distante.

Le protocole **NCP** se décline en autant de versions que de protocoles réseau supportés. On trouve ainsi IPCP (*IP Control Protocol*), DCP (*Decnet Phase IV Control Protocol*), etc. Les trames échangées sont du même type que celles utilisées par LCP.

IPCP (RFC 1332) permet de négocier des options spécifiques, comme l'affectation des adresses IP ou la compression des en-têtes TCP/IP. Pour éviter la fragmentation des paquets TCP, la longueur maximale des données peut être négociée pour aller au-delà des 1 500 octets, et correspondre ainsi au MTU du protocole TCP.

PPP se contente ensuite de véhiculer les paquets IP en offrant simplement la détection d'erreur.

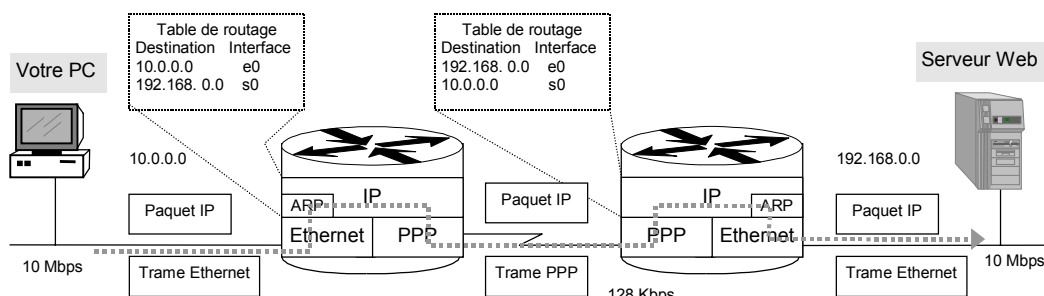
La même configuration doit être appliquée au routeur d'Orléans.

```
int e0
ip address 192.168.0.1 255.255.255.0
int s0
ip unnumbered e 0
encapsulation ppp
^Z
ip route 10.0.0.0 255.255.252.0 s0
```

Grâce à ces commandes, les routeurs sont capables de router correctement les paquets.

Figure 8-5.

*Principe du routage
des paquets IP.*



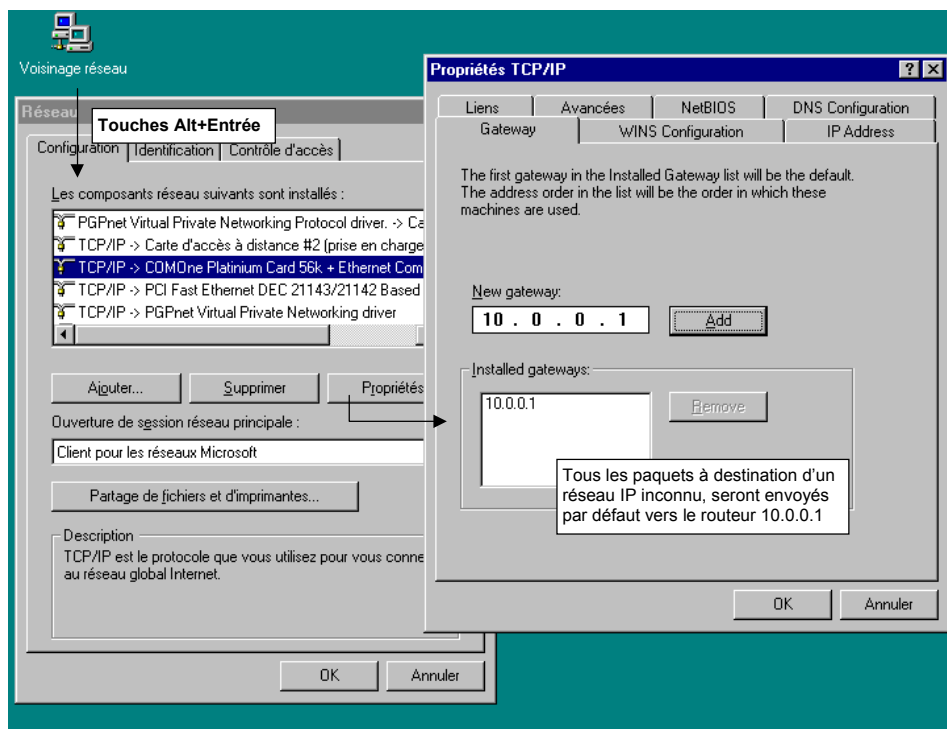
Étant donné qu'il n'existe qu'un seul chemin (la ligne série), il est possible de déclarer une route par défaut. Cette commande permet au routeur d'envoyer sur son interface série tous les paquets dont il ne connaît pas l'adresse de destination :

```
# Remplace la commande ip route 192.168.0.0 255.255.255.0 s0
ip route default s0
```

Configurer les postes de travail

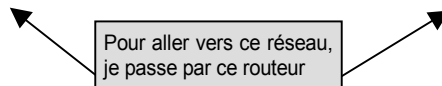
La question est maintenant de savoir comment les PC vont pouvoir envoyer les paquets IP sur l'autre site. Réponse : selon le même principe que celui utilisé par les routeurs.

Il suffit, en effet, d'ajouter une route par défaut (*default gateway*).



Si une route par défaut existe déjà vers un autre routeur, il est toujours possible d'ajouter une route statique, comme suit :

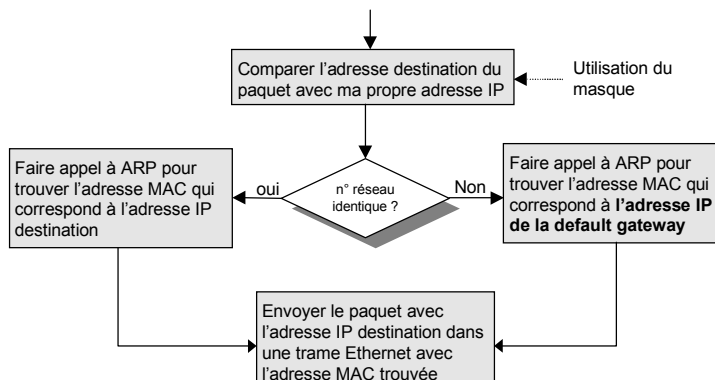
```
route add -p 192.168.0.0 mask 255.255.255.0 10.0.0.1
```



Mais, comment la résolution d'adresse fonctionne-t-elle étant donné que l'adresse MAC du serveur situé à Orléans nous est inconnue et que les trames de broadcast ARP ne peuvent pas passer par le routeur ?

On pourrait activer la fonction Pont du routeur et ne laisser passer que les broadcasts ARP, mais on perdrait alors l'avantage de la segmentation : le réseau de Paris recevrait des broadcasts MAC dont il n'a pas l'usage, et inversement.

En réalité, la solution retenue par la pile IP est la suivante :



En définitive, le PC recherchera l'adresse MAC du routeur et l'associera à l'adresse IP de destination. Les trames Ethernet seront ainsi envoyées directement au routeur. Ce dernier – grâce à l'adresse de destination IP – comprendra que le paquet ne lui est pas destiné. Mais, au lieu d'ignorer le paquet comme le font les PC (voir chapitre 7), il consultera sa table de routage et retransmettra le paquet vers la bonne interface (Ethernet ou autre, ici la série).

Le routeur d'Orléans reçoit le paquet. L'adresse IP de destination correspondant à un réseau auquel il est directement connecté, celui-ci fait appel à ARP pour obtenir l'adresse MAC qui correspond à cette adresse IP. Le routeur envoie alors le paquet IP dans une trame Ethernet dont l'adresse de destination est celle du serveur.



Il est également possible d'indiquer l'**adresse IP de la station** elle-même **comme passerelle par défaut**. Ce faisant, la pile IP considère que tous les subnets lui sont directement accessibles. Par conséquent, si ceux-ci sont tous situés sur le même segment Ethernet, la station pourra directement accéder à toutes les autres stations. Si les subnets se trouvent sur d'autres segments interconnectés par des routeurs, le routeur local doit être configuré en mode **Proxy ARP**, ce qui est le cas par défaut de nos routeurs Cisco.

LE POINT SUR PROXY ARP (RFC 1027)

Le protocole ARP (*Address Resolution Protocol*) permet à une station de connaître l'adresse Ethernet MAC d'une autre station en ne connaissant que son adresse IP. La trame de broadcast envoyée par ARP n'est diffusée que localement à un segment : elle est, en effet, bloquée par les routeurs, comme cela est la règle dans les réseaux locaux.

En mode proxy ARP, un routeur qui reçoit une requête ARP concernant une adresse IP dont il connaît le réseau (car l'adresse est présente dans sa table de routage) s'assure qu'il dispose de la meilleure route, puis répond à la requête en y mettant sa propre adresse MAC. En définitive, **il se substitue à la station cible** qui ne peut pas recevoir une telle requête puisqu'elle se trouve sur un autre segment Ethernet.

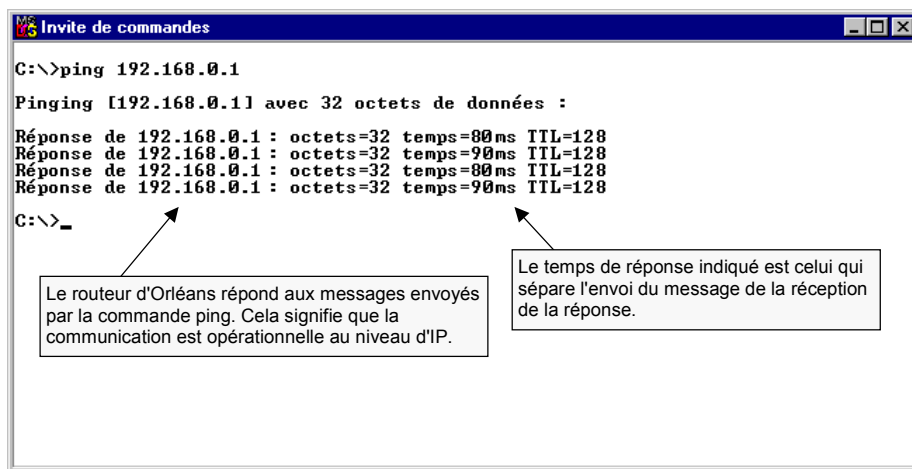
Ces deux mécanismes utilisés conjointement permettent, par exemple, de migrer d'un réseau de routeurs vers un réseau de commutateurs, ou d'un VLAN à plusieurs subnets IP vers un VLAN à un seul. Si, de plus, vous utilisez un serveur DHCP (voir chapitre 16), ce dernier doit également être configuré en conséquence.

Voilà, vous venez de réaliser votre première interconnexion de réseaux.

Tester le réseau

Pour tester le bon fonctionnement de votre réseau, vous pouvez utiliser la commande **ping** (voir chapitre 16). Ce programme est disponible sur les routeurs et sur les PC Windows, les serveur Unix, etc.

Il se contente d'envoyer des paquets à une adresse cible qui lui répond par un paquet en retour. De plus, le ping mesure le temps de réponse aller-retour, ce qui est une bonne indication sur les performances du réseau. Toute pile IP se doit de répondre à un ping.



Optimiser

Une fonctionnalité très souvent utilisée pour les liaisons bas débit (64 à 128 Kbit/s), est la compression des en-têtes TCP/IP (RFC 1144). Le principe repose sur le constat que les en-têtes varient peu d'un paquet à l'autre. Seuls les octets ayant changé par rapport au précédent paquet sont donc transmis. La taille de l'en-tête est ainsi réduite de 40 octets (20 pour TCP et 20 pour IP) à 10 octets en moyenne. Cette fonctionnalité est d'autant plus efficace qu'il y a de nombreux petits paquets à traiter (connexions Telnet, par exemple).

```
interface serial 0
ip tcp header-compression
```

Mettre en place une liaison de secours

Notre liaison est en place : tout fonctionne. Mais que se passera-t-il si elle tombe en panne ? Plus de réseau. La solution consiste à mettre en place une liaison de secours. Là encore, la question du choix du support de transmission et du débit se pose.

Quels sont les choix ?

La première solution consiste à doubler la liaison principale : une seconde LS prend le relais de la première. En fonctionnement normal, les deux liaisons peuvent être utilisées en partage de charge.

L'autre solution consiste à utiliser Numéris, le réseau téléphonique numérique de France Télécom, généralement appelé RNIS (réseau numérique à intégration de service). Pour un coût nettement moins élevé qu'une LS, la liaison RNIS n'est activée que lorsque la liaison principale est coupée.

Secours via...	Avantages	Inconvénients
RNIS	Permet le débordement. Facturé essentiellement à l'utilisation	Limité à 128 Kbit/s en standard et à 384 Kbit/s avec l'agrégation de canaux.
LS	Convient à une utilisation en partage de charge.	Facturée même si elle n'est pas utilisée.

La mise en place de deux LS, par exemple 2×64 Kbit/s à la place d'une seule liaison à 128 Kbit/s, permet de répartir la charge sur les deux liaisons. Cette solution nécessite cependant des mécanismes plus complexes que ceux utilisés jusqu'à présent : ils reposent sur des protocoles de routages qui sont le plus souvent propriétaires. En outre, deux LS à 64 Kbit/s coûtent plus cher qu'une seule à 128 Kbit/s.

Une liaison RNIS présente l'avantage de n'être activée que lorsque cela est nécessaire. Si elle est inutilisée, seul l'abonnement de base doit être payé : de 200 à 300 francs contre 3 000 à 4 000 francs par mois pour une LS 64 Kbit/s courte distance.

Globalement, le secours RNIS revient donc nettement moins cher que le doublement d'une LS. Pour toutes ces raisons, cette solution est donc plus adaptée à notre besoin.

Solutions alternatives

Des techniques d'**agrégation de canaux B** permettent d'utiliser n canaux B à 64 Kbit/s de manière à n'en faire qu'une seule liaison logique à $n \times 64$ Kbit/s. Il est ainsi possible d'agréger les deux canaux B d'un accès de base (offrant un débit global de 128 Kbit/s) ou d'agréger les canaux B de plusieurs accès de base (généralement jusqu'à trois, autorisant un débit global de 384 Kbit/s). Cette technique requiert l'utilisation d'équipements spécifiques qui peuvent être onéreux.



Attention au coût des communications : l'agrégation de deux canaux B offre un débit de 128 Kbit/s mais équivaut à deux communications téléphoniques simultanées. Cette technique coûte donc deux fois plus cher qu'une simple connexion à 64 Kbit/s.

Le choix du débit de la liaison de secours dépend de l'importance que vous accordez à l'interconnexion :

- Si un **mode dégradé** est acceptable, une liaison à 128 ou 256 Kbit/s pourra être secourue à 64 Kbit/s.
- Inversement, si vous voulez conserver les performances en mode secours, il faut utiliser une seconde liaison (LS, satellite ou autre) qui aille au-delà des limitations de débit du RNIS.

LE POINT SUR LE RNIS (ITU SÉRIE I)

Le RNIS (réseau numérique à intégration de service — **ISDN**, en anglais) désigne le réseau téléphonique numérique, par opposition au RTC (réseau téléphonique commuté — **PSTN**, en anglais) qui en est la version analogique. En France, le RNIS est commercialisé par France Télécom sous le nom de Numéris.

Il existe deux types d'abonnements RNIS :

- L'**accès de base T0** qui offre deux canaux à 64 Kbit/s et un canal de signalisation à 16 Kbit/s.
- L'**accès primaire T2** qui revient à fournir un système à trente canaux de 64 Kbit/s chacun et un canal de 64 Kbit/s pour la signalisation (le premier canal, appelé verrouillage de trame, est dédié à la synchronisation du faisceau).

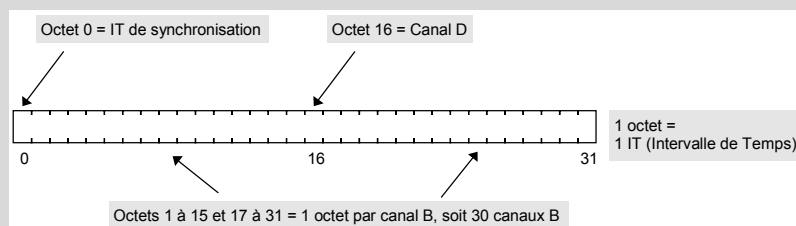
Un canal RNIS permet de transporter une liaison téléphonique, une liaison télécopie ou encore une liaison de données. Cette unité de base est appelée **canal B** (Base), tandis que le **canal de signalisation** est appelé canal D (Données).

L'accès de base est l'équivalent d'un réseau local, appelé **bus S0**, qui est géré par la TNR (terminaison numérique de réseau). Cinq terminaux au maximum (téléphone, télécopieur, routeur ou tout autre équipement pourvu d'une interface de base — BRI, *Basic Rate Interface*) peuvent se partager le bus dont l'accès est géré selon l'algorithme **CSMA-CR** (*Carriage Sense Multiple Access - Contention Resolution*).

Le canal D véhicule le protocole de signalisation de niveau 3 (couche réseau) **Q.931** qui permet d'établir et de gérer les communications entre deux terminaux RNIS (numérotation, identification de l'appelant, négociation des paramètres, etc.).

Les messages Q.931 sont transportés dans des trames de niveau 2 (couche liaison) gérées par le protocole **HDLC** (*High Data Link Control*) qui fournit des mécanismes d'acquittement de trames, de contrôle de flux et de reprise sur erreur pour le canal D. Le RNIS utilise une version adaptée, appelée LAP-D (*Link Access Procedure - D channel*).

La trame de l'accès primaire est constituée de 32 canaux de 64 Kbit/s : 30 canaux B, un canal D et un canal de synchronisation (30B+D, en abrégé). Le débit global de l'accès primaire est de 2 Mbit/s.

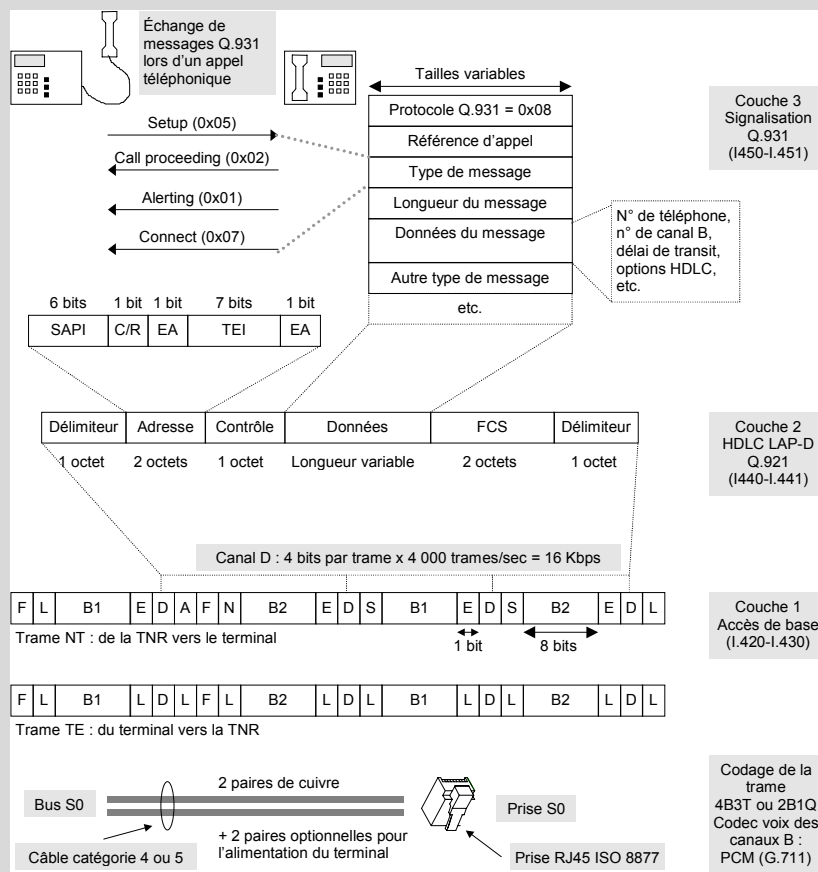


La trame de l'accès de base comprend deux canaux B à 64 Kbit/s, un canal D à 16 Kbit/s et quelques bits de gestion (2B+D, en abrégé). Le débit global est de 192 Kbit/s. D'une longueur de 48 bis, elle est émise en 250 microsecondes (soit 4 000 trames par seconde).

...

LE POINT SUR LE RNIS (SUITE)

L'accès primaire et l'accès de base utilisent les mêmes protocoles de signalisation de niveau 2 et 3.



Délimiteur	est de valeur binaire « 01111110 ». Afin d'éviter que les données lui soient identiques, un mécanisme d' insertion de « 0 » est utilisé à l'émission. Le principe consiste à insérer un « 0 » dès que 5 valeurs « 1 » consécutives ont été transmises (cet algorithme n'est bien sûr pas appliqué aux délimiteurs). À la réception, le 6 ^e bit suivant 5 bits ayant une valeur de « 1 » est analysé. S'il est à « 0 », il est ignoré et, dans le cas contraire, il est considéré comme un délimiteur de trame.
Adresse	Comporte l'identifiant SAPI (<i>Service Access Point Identifier</i>) permettant de distinguer les applications utilisant les canaux B (téléphonie, télécopie, etc.) et l'identifiant TEI (<i>Terminal End-point Identifier</i>) qui permet d'identifier le terminal sur le bus. Le numéro 127 est réservé à la diffusion (<i>broadcast</i>).
Contrôle	indique le type de la trame (<i>Unnumbered</i> , <i>Information</i> ou <i>Supervision</i>), et contient un numéro de séquence. Le protocole utilise 20 formats de trames.
Données	contient les données utiles (260 octets au maximum).
FCS	(<i>Frame Check Sequence</i>) est un code de détection d'erreur de type CRC (<i>Cyclic Redundancy Check</i>).

...

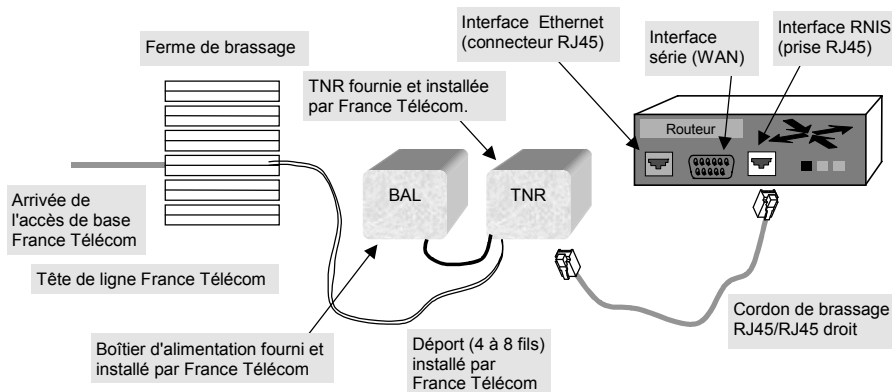
LE POINT SUR LE RNIS (FIN)

F	est un bit de synchronisation.
F _a	est un bit de synchronisation auxiliaire pour une utilisation future.
L	est le bit d'équilibrage qui sert à rétablir la composante énergétique du courant transportant le train de bits. L'objectif est d'obtenir une composante énergétique nulle afin de consommer le minimum d'énergie.
E	bit d'écho utilisé par la méthode d'accès au bus CSMA/CR (<i>Carrier Sense Multiple Access - Contention Resolution</i>).
A	bit utilisé pour l'activation du terminal.
S	bit non utilisé.
B1	est un champ de 8 bits qui véhicule les données du premier canal B.
B2	est un champ de 8 bits qui véhicule les données du deuxième canal B.
D	est le champ qui véhicule les données de signalisation du canal D.

Installation d'un accès de base T0

Comme pour notre LS, l'opérateur (France Télécom en France) installe la liaison RNIS dans la partie de vos locaux techniques qui lui est réservée. L'accès de base T0 se termine chez l'utilisateur par un petit boîtier appelé **TNR** (terminaison numérique de réseau). Celui-ci est, la plupart du temps, accompagné d'un coffret d'alimentation appelé BAL.

Figure 8-6.
Connexion
du routeur
au RNIS.



Sur notre routeur d'Orléans, l'interface RNIS (appelée *bri* pour *Basic Rate Interface*) doit être configurée de manière à ne s'activer que lorsque la liaison principale (l'interface série) est coupée :

```
isdn switch-type vn3
interface Ethernet 0
ip address 192.168.0.1 255.255.255.0
^Z
```

Interface RNIS de France Télécom
(Version Numéris 3)

```

interface serial 0
 ip unnumbered e0
 encapsulation ppp
 backup delay 2 15

backup interface bri 0
^Z
interface bri 0
 ip unnumbered e0
 encapsulation ppp
 dialer string 0144758899
 dialer-group 1
^Z
 ip route 10.0.0.0 255.255.252.0 s0
 dialer-list 1 protocol ip permit

```

Active l'interface RNIS 2 secondes après la chute de l'interface principale. Attend 15 secondes avant de couper la liaison RNIS lorsque l'interface principale est de nouveau opérationnelle.

L'interface de secours est la bri 0 qui appelle le numéro de téléphone de l'abonnement RNIS situé à Paris.

Définition d'une route statique.

En mode secours, un seul canal B est activé. Notre routeur est capable de gérer deux connexions simultanées, une sur chaque canal B, mais pas d'agréger les canaux.

Sécurisation de la liaison

Le problème avec le numéro de téléphone est qu'il peut être appelé par n'importe qui. Il faut donc appliquer quelques règles de base. La première est, bien sûr, de demander que votre numéro soit inscrit sur liste rouge. Le service d'identification de l'appelant pourra aussi être utile.

Au niveau du routeur, il est possible de réaliser une authentification à l'aide du protocole CHAP (*Challenge Handshake Authentication Protocol*) utilisé conjointement avec PPP. Le principe repose sur le partage d'une clé secrète par les deux routeurs. Lors de l'appel, le routeur appelé envoie un « challenge » au routeur appelant. La bonne réponse est liée à la clé secrète (le mot de passe).

```

hostname orleans

username paris password 7 14041BBEAB04

isdn switch-type vn3

interface Ethernet 0
 ip address 192.168.0.1 255.255.255.0

```

Nom du routeur distant qui correspond au "hostname paris".

Clé secrète CHAP (mot de passe)

Le chiffre 7 indique que le mot de passe n'est pas affiché en clair. Attention cependant : il faut saisir le mot de passe en clair lors de la configuration.

```

^Z
interface serial 0
 ip unnumbered e0
 encapsulation ppp
 backup delay 2 15
 backup interface bri 0
^Z
interface bri 0
 ip unnumbered e0
 encapsulation ppp
 dialer in-band
 dialer map ip 10.0.0.0 name paris broadcast 0144758899
 dialer-group 1
^Z
ip route 10.0.0.0 255.255.252.0 s0
dialer-list 1 protocol ip permit

ip route 10.0.0.0 255.255.252.0 s0

```

Pour joindre le réseau 10.0.0.0, il faut appeler le routeur distant "paris".

Numéro de téléphone RNIS du site parisien

Seul le protocole IP est permis.

Gestion du débordement

La même liaison RNIS peut également être utilisée pour absorber du trafic en surplus sur l'interface série :

```

int bri 0
backup load 80 10

```

Cette commande active l'interface RNIS si la ligne principale atteint une charge égale à 80 %, et la déconnecte lorsque la charge globale (principale + secours) redescend à 10 %.

Le cas de la panne d'un routeur n'est pas traité ici. Pour offrir cette sécurité supplémentaire, d'autres mécanismes doivent être activés (voir chapitre 11).

Architecture des réseaux étendus

Après avoir réalisé notre première interconnexion, nous pouvons passer à l'étape suivante qui consiste à créer un réseau WAN complet.

À l'inverse d'un LAN qui est par essence privé, un WAN nécessite d'emprunter des réseaux publics ou opérateurs qui agissent sous licence octroyée par l'État.

À l'inverse de votre LAN sur lequel vous pouvez réaliser des excès de vitesse gratuitement jusqu'au gigabit, la vitesse est limitée à quelques dizaines de mégabits sur les réseaux WAN. Et, plus vous allez vite et loin, plus c'est cher.

Il faut donc choisir avec discernement la technologie à utiliser, et déterminer au mieux en fonction de vos besoins et de votre budget la route à emprunter (route départementale, voie rapide, autoroute de l'information avec péage, bretelle d'accès, etc.).

Dans ce chapitre, vous apprendrez ainsi :

- à choisir un niveau de service opérateur ;
- à connaître les technologies d'accès xDSL ;
- à estimer la volumétrie générée par vos applications ;
- à dimensionner votre réseau WAN.

Les solutions disponibles sur le marché

Après notre première interconnexion, voilà qu'il faut interconnecter quatre nouveaux sites : Toulouse, Marseille, Strasbourg et Londres.

Quels sont alors les moyens mis à notre disposition ?

Les infrastructures

L'élément de base d'un réseau étendu est la **liaison** détenue et gérée par un **opérateur** : France Télécom et Cegetel en France, Deutsche Telekom en Allemagne, MCI et Sprint (pour ne citer qu'eux) aux États-Unis, etc.

Une liaison peut être **filaire** (cuivre ou fibre optique) ou **hertzienn**e (satellites ou émetteurs terrestres). Les liaisons filaires sont obligatoirement **point à point**, tandis que les liaisons hertziennes peuvent être point à point ou **multipoint**.

Les opérateurs se louent leurs liaisons entre eux. Ils bâtissent leurs propres réseaux basés sur des liaisons qui leur appartiennent et sur d'autres qu'ils louent là où ils ne peuvent pas en construire.

La construction de liaisons filaires ou hertziennes requiert, en effet, des autorisations administratives (**licences**) pour pouvoir poser les câbles à travers des territoires appartenant à chaque État, lancer les satellites, obtenir les fréquences hertziennes et les exploiter commercialement. En France, l'ART (autorité de régulation des télécommunications) est responsable de l'attribution de ces licences aux opérateurs. Depuis peu, la libéralisation du commerce mondial incite à faire sauter les derniers monopoles.

Au niveau international, l'organisme qui contrôle les activités télécoms (autorisations, projets internationaux, normes, etc.) est l'**ITU** (*International Telecommunication Union*) qui est affilié à l'ONU. Tous les opérateurs nationaux sont membres de l'ITU.

La plupart des liaisons internationales sont le fruit d'une coopération entre les opérateurs qui utilisent les bandes passantes proportionnellement à leur participation financière. Certains opérateurs possèdent en propre leurs liaisons internationales, et les louent à d'autres opérateurs.

RÉSEAUX LOCAUX, ÉTENDUS ET INTERSITES

Le terme LAN (*Local Area Network*) désigne les **réseaux locaux**, dont le principal représentant est Ethernet. Le terme WAN (*Wide Area Network*) désigne les **réseaux étendus** dont les représentants les plus répandus sont les LS, Frame Relay et ATM.

On désignera par **réseau intersite** un réseau d'interconnexion de réseaux locaux reposant sur un réseau étendu et des **routeurs** qui réalisent l'interface entre les LAN et le WAN.

Les réseaux opérateurs

À partir des liaisons qui lui appartiennent en propre ou qu'il loue, un opérateur crée un ou plusieurs réseaux interconnectant ses sites. Ces liaisons se terminent par des multiplexeurs et des commutateurs (Frame-Relay, ATM, SDH ou propriétaires). L'opérateur propose ensuite à ses clients de partager son réseau en leur revendant de la bande passante et en leur proposant un service d'exploitation.

Une entreprise désirant interconnecter ses sites, ou un particulier désirant communiquer à distance (téléphoner, se connecter à l'Internet ou à son intranet, etc.) devra **obligatoirement faire appel aux services d'un opérateur**.

Le premier de ces réseaux est historiquement le **RTC** (ou réseau téléphonique commuté) qui permet aux entreprises et aux particuliers de communiquer par téléphone moyennant un abonnement et une facturation à la durée.

Le second type qui est apparu est la **LS** (ligne spécialisée, encore appelée ligne louée), qui permet aux entreprises d'interconnecter leurs sites moyennant un abonnement mensuel dont le coût dépend du débit utilisé et de la distance.

L'accès au réseau

Les sites et les liaisons qui constituent le réseau de l'opérateur peuvent être plus ou moins nombreux, et la **zone de couverture** de ce réseau peut être plus ou moins étendue.

Chaque client doit se raccorder au réseau de l'opérateur *via* des **liaisons d'accès** (filaires le plus souvent, ou hertzienne pour le téléphone mobile, par exemple) entre son site et les points d'accès appelés **POP** (*Point Of Presence*) de l'opérateur. La liaison d'accès est également appelée desserte locale ou, plus spécifiquement, **boucle locale**.

L'intérêt d'une telle solution est que le POP soit plus près possible du site client afin que la liaison de raccordement, dont le prix dépend du débit et de la distance, coûte le moins cher possible. Le RTC est le réseau opérateur qui dispose du plus important nombre de POP : il s'agit d'un commutateur installé au coin de la rue, dans un immeuble. Mais, pour les autres réseaux, il n'en est pas de même. Par exemple, vous avez sans doute consulté les zones de couverture d'Itinéris, de Bouygues Telecom, de SFR, etc.

L'opérateur installe toujours un équipement d'extrémité, appelé **CPE** (*Customer Premises Equipment*), dans les locaux du client qui permet de gérer la liaison d'accès. Il s'agit, par exemple, d'un modem, d'un commutateur, etc., selon le niveau de service fourni.

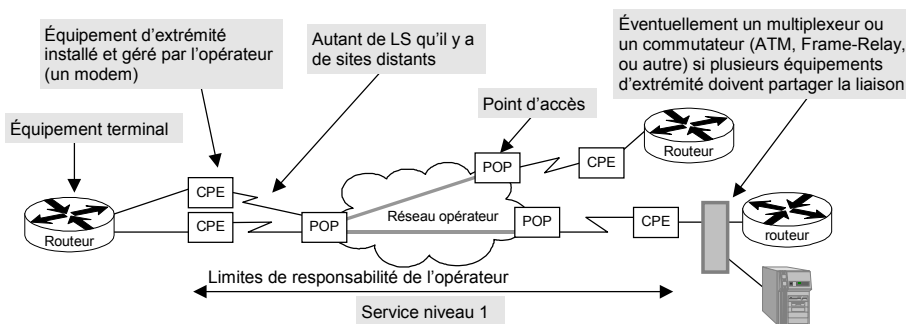
Les services proposés par les opérateurs

L'opérateur propose à ses clients d'utiliser son réseau en leur revendant de la bande passante ainsi que différents niveaux de service d'exploitation. À la diversité des technologies s'ajoute le maquis des services proposés, l'habillage commercial en quelque sorte. Bien qu'il en existe de nombreuses variantes, on trouve principalement trois types de services.

Niveau de prestation	Description technique	Service fourni	Exemples
1. Fourniture du support de transmission brut (couche physique)	Liaisons (le plus souvent point à point) pour créer un réseau privé	Support client de la liaison	Lignes spécialisées, liaisons VSAT
2. Réseau fédérateur (couches physique et liaison)	Création d'un réseau privé (le plus souvent multipoint) au-dessus de l'infrastructure opérateur	Réseau fourni et exploité par l'opérateur + support client	Réseau ATM, Frame-Relay ; accès à l'Internet
3. Service à valeur ajouté (couches physique, liaison et réseau)	Support de transmission + réseau fédérateur + équipements terminaux (routeur, téléphone, etc.)	Réseau de bout en bout (juste dans le site du client) fourni et exploité par l'opérateur	RTC, RNIS, interconnexion de LAN, accès à l'Internet

Le service de niveau 1 permet aux clients de disposer du support de transmission qui leur est fourni : ils peuvent installer des multiplexeurs, des commutateurs ATM ou Frame-Relay, des routeurs, etc. pour transmettre des données, de la voix ou de la vidéo.

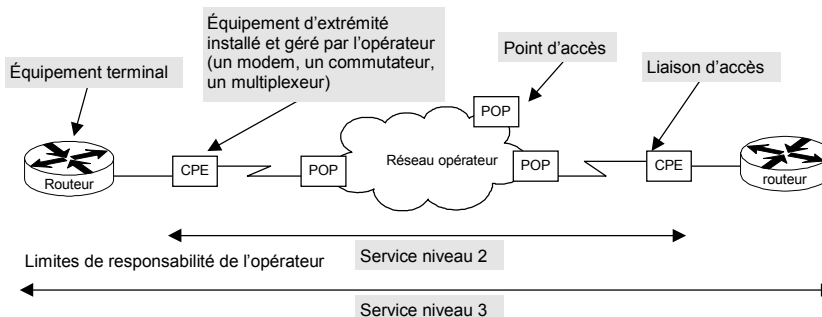
Figure 9-1.
Service opérateur
de niveau 1
– lignes
spécialisées.



Le service de niveau 2 consiste à profiter de l'infrastructure réseau de l'opérateur, qui est vu comme un nuage (un réseau multipoint), sur lequel les sites du client sont raccordés *via* des liaisons locales et les POP.

Le service de niveau 3 est un service complet : l'opérateur propose un service clé en main *via* un réseau quelconque, avec les moyens et les technologies qu'il souhaite. Le client ne voit que le service : téléphone, interconnexion de réseaux locaux, etc.

Figure 9-2.
Services opérateurs
de niveaux 2 et 3.



Pour compliquer un peu la situation, sachez qu'il existe deux types d'offres de niveaux 2 et 3 :

- Les **VPN** (*Virtual Private Network*). Les clients se partagent le réseau de l'opérateur mais sont physiquement ou logiquement séparés. L'opérateur leur garanti la bande passante demandée ainsi que le niveau de service souhaité, avec des engagements de résultats.
- Les **VPN-IP** (*VPN sur un réseau IP*). Les clients se partagent le réseau de l'opérateur mais ne sont pas obligatoirement séparés. L'opérateur ne garantit pas que le client disposera de toute la bande passante demandée et n'offre généralement pas ou peu d'engagements de services.

Les VPN sont bien sûr plus chers que les VPN-IP, car les engagements de services sont meilleurs. Les VPN sont pour cela adaptés aux interconnexions de réseaux locaux, tandis que la vocation première des VPN-IP est la connexion à l'Internet, voire la mise en place d'un intranet.

Dans tous les cas, le client doit signer un contrat de service, appelé **SLA** (*Service Level Agreement*), avec l'opérateur. Ce dernier s'engage à fournir une qualité de service (temps de réponse, débit garanti, taux de disponibilité, etc.) ainsi qu'à payer des pénalités au client en cas de non-respect de ses engagements.

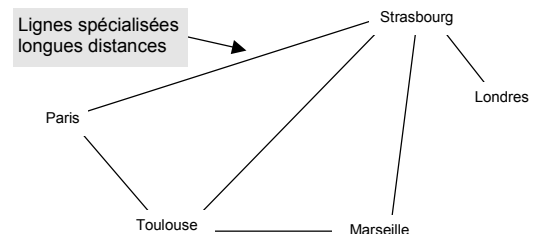
Les choix du client

Pour créer notre réseau étendu, la première question concerne le niveau de service que nous allons demander à l'opérateur. Du point de vue client, nous avons le choix entre deux types de solutions :

- **Solution privée** reposant sur le service de niveau 1. L'entreprise construit son réseau étendu.
- **Solution opérateur** reposant sur les services de niveau 2 et 3. L'entreprise confie tout ou partie de son réseau étendu à un opérateur.

La première solution consiste à suivre le modèle de notre première interconnexion, afin de réaliser notre propre réseau multipoint reposant sur des lignes spécialisées point à point.

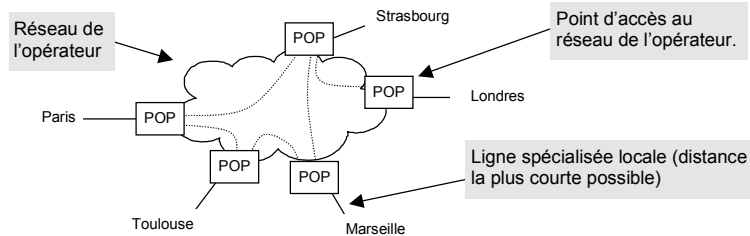
Figure 9-3.
Réseau privé reposant sur des LS.



Le réseau peut être en étoile sur un site principal, distribué selon l'importance relative de chaque site, ou plus ou moins maillé afin de répartir la charge et d'assurer le secours des liaisons.

La seconde solution repose sur l'utilisation d'un réseau opérateur. Elle consiste en des liaisons d'accès aux POP que l'on peut considérer comme étant des lignes spécialisées locales. Pour une interconnexion de réseaux locaux, l'offre VPN est plus appropriée que l'équivalent sur IP.

Figure 9-4.
*Réseau privé virtuel
basé sur un réseau
opérateur.*



L'opérateur installe les liaisons d'accès locales (ou les commande auprès de l'opérateur local s'il n'a pas les licences) entre les sites du client et ses points d'accès.

Le réseau opérateur peut s'engager sur différents services :

- Taux de disponibilité. L'opérateur garantit que son réseau sera disponible 99,9 % du temps (c'est une valeur courante).
- Bande passante. L'opérateur garantit que le client disposera du débit demandé (512 Kbit/s entre deux sites, par exemple) pendant 100 % du temps.
- Temps de transit. L'opérateur garantit le temps mis par un paquet pour aller d'un site à l'autre.

Ces garanties peuvent être valables de POP à POP (c'est-à-dire dans le réseau fédérateur de l'opérateur appelé backbone) ou de bout en bout (c'est-à-dire entre les sites du client), liaisons et équipements d'accès compris. C'est à vous de négocier les engagements en fonction du coût.

L'opérateur exploite son réseau à l'aide d'un ou plusieurs centres de supervision, de centres de support client (*Help Desk*) et d'équipes projet qui lui permettent d'offrir un service de guichet unique (interlocuteur unique) :

- L'opérateur se charge de toutes les commandes de LS d'accès auprès des opérateurs locaux, de l'installation des équipements, de la configuration de son réseau pour accueillir le VPN du client ainsi que de la gestion du projet.
- Quels que soient le site et le pays, le client peut appeler le centre de supervision pour lui signaler un problème, et inversement.
- Il n'y a qu'une facture, et le client choisit le mode de facturation qu'il souhaite : centralisée, répartie par site, etc.

- L'opérateur fournit des statistiques sur le VPN du client et les preuves que ses engagements ont été respectés. Cela n'empêche pas le client de mettre en place son propre système de contrôle.

Critère	Solution privée	Solution opérateur
Investissements	Importants (frais de mise en service et achat des équipements)	Faibles (frais de mise en service) ; pas d'achat d'équipement
Coûts de fonctionnement	Assez faibles en national, élevés à l'international	Assez élevés (coût des LS plus faible et forfait opérateur)
Exploitation	Assurée par la société (vous)	Assurée par l'opérateur (supervision des LS, exploitation, etc.)
Distances des LS	Longues entre les sites	Locales entre les sites et les POP de l'opérateur

Une entreprise a toujours la possibilité de passer d'une solution privée à une solution opérateur, autrement dit d'**externaliser** son réseau (ce que l'on appelle l'**outsourcing**). L'opération inverse (la **réversibilité**) n'est pas dans l'air du temps et est toujours une opération délicate. La réversibilité peut aussi être appliquée pour changer d'opérateur en cas d'insatisfaction.

Le réseau de transport

Les opérateurs utilisent des réseaux spécifiques pour transporter une variété de flux (voix, lignes spécialisées, Frame-Relay, ATM, etc.) à des débits divers (de 64 Kbit/s à plusieurs Mbit/s). Les opérateurs doivent mettre en place une infrastructure pour constituer leur réseau fédérateur (*backbone*) et pour en permettre l'accès à leurs clients (boucle locale).

Qu'est-ce qu'une LS ?

Une LS (ligne spécialisée ou ligne louée) est un terme générique pour désigner une liaison point à point entre deux sites.

Pour l'utilisateur, une LS consiste en deux ou quatre fils (plus rarement six) de cuivre reliant deux de ses sites. Pour l'opérateur, ce n'est qu'une liaison d'accès à son réseau : la liaison cuivre qui part du site du client aboutit à un multiplexeur qui concentre d'autres LS (issues d'autres clients) sur des liaisons haut débit en fibre optique.

Les LS reposent sur deux technologies :

- T1 (États-Unis, Canada et Japon) et E1 (reste du monde), qui date des années 60.
- HDSL (*High bit rate Digital Subscriber Line*), qui date des années 80.

Dénomination	Technologie	Codage en ligne	Distance
E1	MIC (<i>modulation par impulsions codées</i>) 32 canaux de 64 Kbit/s	Bipolar AMI (<i>Alternate Mark inversion</i>)	De 1 à 2 km sans répéteur 2 paires
T1	PCM (<i>Pulse Code Modulation</i>) 24 canaux de 64 Kbit/s	Codage bipolar AMI	De 1 à 2 km sans répéteur 2 paires
HDSL	DSL Trames transmises en 6 ms	2B1Q (2 bits / 1 signal quaternaire)	De 3,7 à 7,9 km sans répéteur 1, 2 ou 3 paires

Le débit offert par une E1/T1 est un multiple de 64 Kbit/s. Cette unité correspond à un canal dans un lien E1 à 2,048 Mbit/s (32 canaux à 64 Kbit/s) ou T1 (*Trunk-carrier, level 1 multiplexing*) à 1,544 Mbit/s (24 canaux à 64 Kbit/s + 8 Kbit/s de signalisation).

Les ondes sonores de la voix sont converties en signaux numériques à un débit de 64 Kbit/s, tout comme le mouvement est converti en 24 images/seconde par une caméra. Ce débit est lié aux limitations technologiques des années 60, et demeure encore l'unité de référence. Les progrès font que, aujourd'hui, on peut se contenter de 8 Kbit/s, voire moins.

L'unité de base est donc constituée d'un **canal** de 64 Kbit/s (dénommé DS0, *Digital Signaling 0*). Étant donné qu'il est plus pratique et plus économique de transporter plusieurs canaux en même temps, ces derniers sont multiplexés au sein d'une liaison composite (*trunk*). Et, comme dans toutes les couches réseaux (nous sommes ici au niveau physique), un protocole est nécessaire pour gérer ces canaux (début et fin du canal, synchronisation des horloges, etc.). Un canal est donc dédié à ce protocole souvent appelé canal de **signalisation**.

Liaison	Canaux de données utiles	Canaux de signalisation
RNIS (réseau téléphonique numérique)	2 canaux à 64 Kbit/s	1 canal à 16 Kbit/s
E1 en Europe : 2 048 Kbit/s	30 canaux à 64 Kbit/s	2 canaux à 64 Kbit/s
T1 aux États-Unis : 1 544 Kbit/s	24 canaux à 64 Kbit/s	1 canal à 8 Kbit/s

Ces échanges point à point sont réalisés entre deux multiplexeurs qui dialoguent *via* le canal de signalisation (des bits prélevés sur le débit global).

Un multiplexeur prend ainsi n canaux à 64 Kbit/s en entrée, et génère un signal de 2 048 Mbit/s en sortie pour une E1. Un autre type de multiplexeur prend quatre canaux à 2 Mbit/s en entrée, et génère un signal à 8 Mbit/s en sortie, et ainsi de suite. On définit ainsi une hiérarchie de débits **plésiochrones** (plusieurs horloges, une pour chaque type de multiplexeur) jusqu'à 34 Mbit/s en Europe et 45 Mbit/s aux États-Unis.

Aujourd'hui, cette cascade de multiplexage a été remplacée par des multiplexeurs permettant d'extraire directement la bande passante souhaitée. On définit ainsi une hiérarchie de débits **synchrones** (une seule horloge pour transporter plusieurs débits). La structure des trames et le protocole associé sont **SONET** (*Synchronous Optical Network*) aux États-Unis, et **SDH**

(*Synchronous Digital Hierarchy*) en Europe. Ces réseaux forment le cœur des infrastructures réseau haut débit des opérateurs (de 51,84 Mbit/s à plus de 9 Gbit/s).

L'opérateur installe dans les locaux du client un **CSU/DSU** (*Channel Service Unit/Data Service Unit*) qui génère et reçoit les signaux sur la LS. Le CSU/DSU est l'équivalent numérique du modem : alors que ce dernier convertit les signaux analogiques en signaux numériques et inversement, le CSU/DSU convertit les signaux numériques des interfaces locales (V.35, X21/V11, etc.) en signaux numériques adaptés aux longues distances, xDSL par exemple.

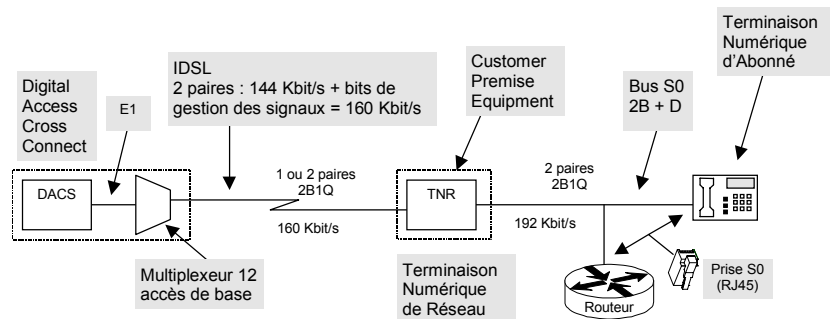
Le **CPE** (*Customer Premises Equipment*) est le terme générique désignant un équipement de raccordement installé chez le client. Il s'agit d'un CSU/DSU (équipement de base), d'un commutateur ATM ou Frame-Relay, d'un FRAD ou d'un routeur. Ces derniers peuvent ou non intégrer un CSU/DSU.

La boucle locale

Les technologies PCM/MIC utilisées par les liaisons T1 et E1 depuis les années 60 sont aujourd'hui dépassées. Les lignes spécialisées offrent toujours une interface E1 ou T1 à leurs clients, mais le PCM a fait place aux technologies DSL (*Digital Subscriber Line*). Celles-ci utilisent des codages plus performants et des processeurs spécialisés dans le traitement du signal, les DSP (*Digital Signaling Processing*).

La première de ces technologies est **IDSL** (*Integrated Service Digital Network DSL*), couramment appelée RNIS (*réseau numérique à intégration de service*) et commercialisée sous le nom de Numéris par France Télécom. Le RNIS est la base du réseau téléphonique numérique ; il a été la première technologie numérique à être proposée aux particuliers. L'accès de base T0 offre deux canaux B de 64 Kbit/s chacun pouvant être utilisés séparément (deux communications) ou agrégés pour offrir un débit de 128 Kbit/s. Un canal D de 16 Kbit/s est réservé à la signalisation (numérotation, réveil du terminal, etc.). Des bits supplémentaires permettent de gérer les signaux transmis sur le câble.

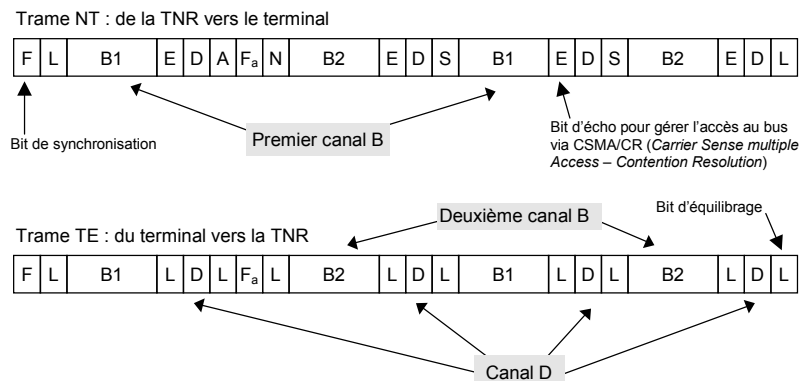
Figure 9-5.
L'accès
de base IDSL.



Côté client, le RNIS se présente sous la forme d'un bus, appelé bus S0, sur lequel plusieurs équipements peuvent être connectés (téléphones, télécopieurs, routeurs, etc.).

Figure 9-6.

Les trames IDSL
(accès de base S0).



La seconde technologie qui est apparue est **HDSL (High bit rate DSL)** sur laquelle reposent quasiment toutes les liaisons d'accès (LS notamment). La raison est qu'une liaison E1/T1 nécessite de coûteux répéteurs tous les 1 à 2 km, alors que la portée du HDSL est d'au moins 3,7 km (jusqu'à 7,9 km sans répéteur). La réduction du nombre de répéteurs permet ainsi de réduire les coûts des LS de 30 à 50 % par rapport aux E1/T1.

La vraie révolution des technologies xDSL est que cette baisse de prix permet aux opérateurs de proposer des liaisons xDSL aux particuliers pour leurs connexions téléphonique et Internet.

Figure 9-7.

Accès E1
via HDSL.

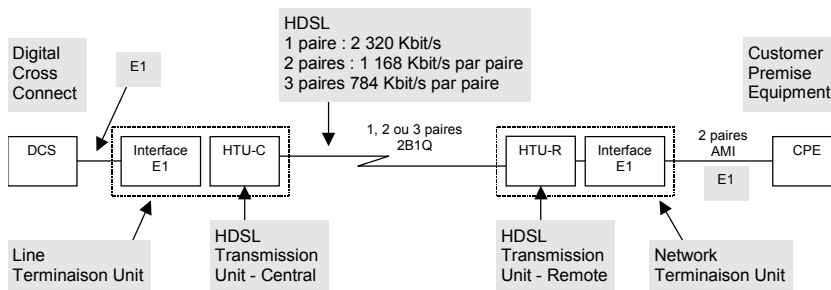
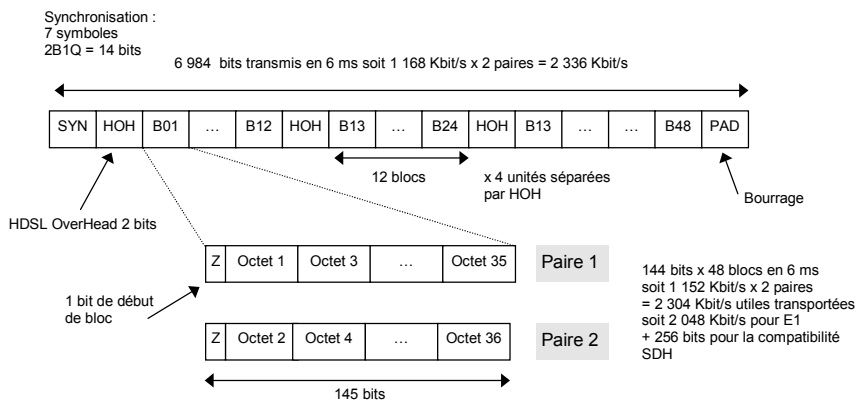


Figure 9-8.

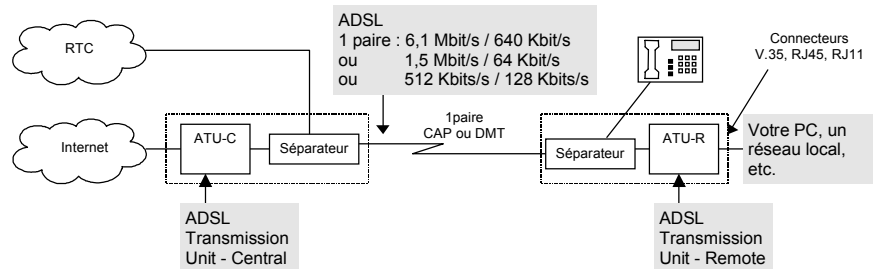
Trames HDSL
sur deux paires.



S-HDSL (*Single pair HDSL*) et **SDSL** (*Symetric DSL*) sont des versions simplifiées de HDSL.

La norme DSL la plus avancée est **ADSL** (*Asymetric DSL*). Sa particularité est de présenter des débits différents selon le sens de la transmission : le débit à destination du client est plus important que celui offert à ce dernier en émission. Ce type d'accès est associé à un séparateur offrant un accès analogique pour les téléphones classiques en plus de l'accès numérique. Cela en fait une utilisation appropriée pour les particuliers.

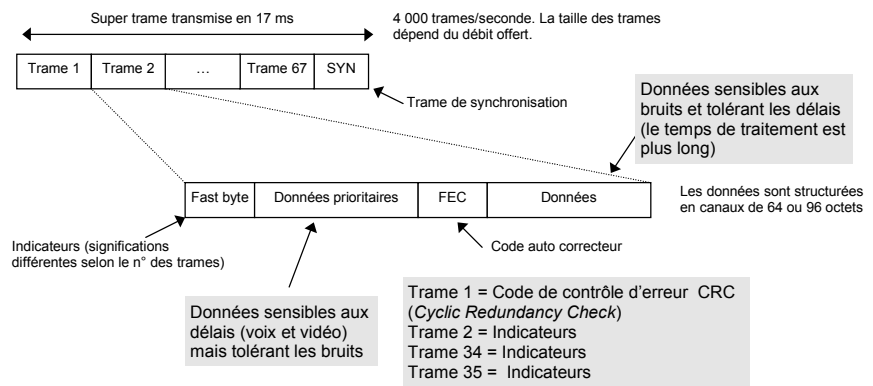
Figure 9-9.
*Accès ADSL
pour
les particuliers.*



Le séparateur peut être inclus ou non dans l'ATU.

Il est à noter que les TNR, HTU et ATU qui viennent d'être évoqués aux figures précédentes sont des modems numériques (CSU/DSU).

Figure 9-10.
*Format
d'une trame
ADSL.*



RADSL (*Rate Adaptive DSL*) est une version d'ADSL qui permet d'augmenter ou de diminuer le débit selon ce que permet la qualité de la ligne. C'est en fait une évolution naturelle d'ADSL.

CDSL (*Customer DSL*) est l'équivalent d'ADSL sans le séparateur (ou séparateur intégré, le draft n'est pas clair à ce sujet). L'objectif est de diminuer les coûts en simplifiant l'installation de l'abonné (suppression du séparateur) ainsi que l'exploitation pour l'opérateur. Les fonctionnalités sont les mêmes que celles fournies par ADSL/RADSL (connexion Internet + téléphone).

La dernière technologie en date est le **VDSL** (*Very high speed DSL*) qui utilise les paires torsadées en cuivre mais surtout la fibre optique. L'objectif premier de VDSL est de transporter des cellules ATM (*Asynchronous Transfer Mode*). La norme repose sur un multiplexage temporel (TDM, *Time Division Multiplexing*) aux débits normalisés STM (*Synchronous Transfer Mode*), ceux utilisés par SDH (*Synchronous Digital Hierarchy*).

Technologie	Débit descendant / Montant en Kbit/s	Distance maximale sans répéteur	Nombre de paires	Codage
HDSL <i>High bit rate DSL</i>	De 1 544 / 1 544 à 2 048 / 2 048	3,7 km (24 AWG) 2,7 km (26 AWG)	2 1 à 3	2B1Q
ADSL <i>Asymmetric DSL</i>	1 544 De 6 144 / 640 à 1 544 / 176	5,5 km 3,7 km 5,5 km	1 1 1	CAP CAP/DMT DMT
RADSL <i>Rate Adaptive DSL</i>	De 1 544 / 64 à 6 144 / 640	5,5 km 3,7 km	1	DMT
CDSL <i>Consumer DSL</i>	1 024 / 128	5,5 km	1	CAP/DMT
VDSL <i>Very high speed DSL</i>	De 13 000 / 1 500 à 52 000 / 6 000	1,3 km 0,304 km	1 ou FO	DWMT / SLC
S-HDSL <i>Symmetric HDSL</i>	768 / 768	3,7 km	1	2B1Q
SDSL <i>Single pair HDSL</i>	De 128 / 128 à 1 024 / 1 024	De 3,5 km à 6,7 km	1	
IDSL <i>ISDN DSL</i>	160 / 160		2	2B1Q

Le débit descendant correspond aux flux allant du réseau (par exemple l'Internet) vers le client (vous), tandis que le débit montant correspond au flux allant du client vers le réseau.

Au sein d'une même norme, les débits et distances varient en fonction du diamètre des fils utilisés (norme AWG, *American Wire Gauge*).

Diamètre des fils	Distance maximale
22 AWG = 0,63 mm	7,9 km (26 000 pieds)
24 AWG = 0,5 mm	5,5 km (18 000 pieds)
26 AWG = 0,4 mm	3,7 km (12 000 pieds)

Les débits peuvent, par ailleurs, être augmentés si les distances sont raccourcies.

Codage utilisé par xDSL	Brève description
DMT (Discrete Multi-Tone)	Repose sur les transformations de Fourier pour gérer et démoduler 256 sous-canaux (sous-porteuses).
DWMT (Discrete Wavelet Multi-Tone)	Repose sur une fonction mathématique, les ondelettes, plus performante que les transformations du Fourier.
SLC (Simple Line Code)	Codage en bande de base à quatre niveaux.
QAM – 16 (Quadrature Amplitude Modulation)	Deux amplitudes et douze changements de phase permettent d'obtenir seize signaux différents représentant 4 bits de données.
CAP (Carrierless Amplitude Phase)	Analogue à QAM, mais sans générer de porteuse.
PAM (Pulse Amplitude Modulation)	2B1Q est un exemple de code PAM à quatre niveaux : 2 bits de données sont codés en un signal quaternaire (quatre niveaux électriques).

Les applications xDSL

Suite à la libéralisation du marché des télécoms, l'enjeu commercial qui aiguise le plus les appétits est celui de la **boucle locale**, c'est-à-dire la liaison d'accès aux réseaux des opérateurs. HDSL couvre les besoins des entreprises en matière de réseau étendu, tandis que les besoins des particuliers (téléphone, télévision, accès à l'Internet – c'est-à-dire les services résidentiels) sont couverts par ADSL/RADSL/CDSL.

Technologie	Applications / Marché visé
IDSL (accès de base RNIS, 2B+D)	Pour les entreprises : Téléphonie numérique, accès à l'Internet et à l'intranet Interconnexion des réseaux locaux : accès principal si les temps d'utilisation sont faibles, secours d'une liaison principale et débordement en cas de surcharge de la ligne principale
HDSL	Accès primaire RNIS (2 Mbit/s) Liaison d'accès E1/T1 de 64 Kbit/s à 2 Mbit/s
ADSL, RADSL, CDSL	Services résidentiels : téléphone, télévision, vidéo à la demande, connexion à l'Internet
VDSL	Réseaux haut débit ATM sur cuivre et surtout sur fibre optique

Concernant nos réseaux locaux, PPP fonctionne directement au-dessus des LS (donc de HDSL) et de RNIS (IDSL) grâce aux routeurs qui prennent en compte les interfaces E1/T1 et S0.

Quant à ADSL, quatre modes d'accès aux canaux de données sont possibles :

- Mode **synchrone**. Canaux accessibles sous forme d'un train de bits aux débits STM.
- Mode **adaptation de paquet**. Permet à plusieurs applications d'utiliser les canaux pour transporter leurs données selon leurs propres formats.
- Mode **paquet de bout en bout**. Les paquets IP sont envoyés directement dans les trames (sans se soucier de l'affection des canaux), et la commutation est effectuée au niveau d'ADSL sur la base des adresses IP contenues dans les paquets.
- Mode **ATM**. Permet de transporter les cellules ATM qui contiennent les trames PPP (qui transportent IP, et ainsi de suite)

Actuellement, aucune tendance n'est perceptible, car le marché est naissant et les équipements peu nombreux. En France, France Télécom promeut le mode ATM.

Dimensionner les liaisons

Le choix du débit est important, car il influe directement sur le coût des liaisons. Celui-ci est d'ailleurs d'autant plus élevé que la distance entre les deux sites est grande. Par exemple, une simple liaison 64 Kbit/s entre la France et la Chine coûte plusieurs dizaines de milliers de francs par mois. S'il s'agit d'une liaison locale vers un POP de l'opérateur, le coût est moindre, mais ce dernier facturera de toute façon en fonction du débit.

Il ne faut donc pas surévaluer le débit par rapport à nos besoins, afin d'éviter de payer un surcoût inutile. Il ne faut pas non plus le sous-évaluer, car les utilisateurs exigent des temps de réponse corrects. La conception d'un **réseau intersite** (réseau d'interconnexion de réseaux locaux) résulte donc d'un compromis coûts/performances.

La démarche proposée pour dimensionner les liens repose sur trois étapes :

- Identifier les flux générés par les applications.
- Estimer la volumétrie, soit à partir de statistiques existantes (facturation, traces relevées sur les équipements, etc.), soit à partir d'hypothèses.
- Déterminer une formule permettant de calculer le débit nécessaire.

Identifier les flux

Nos utilisateurs sont répartis sur six sites (le siège — Paris — et les directions régionales : Orléans, Toulouse, Marseille, Strasbourg et Londres). Ils veulent utiliser les mêmes applications et accéder aux mêmes données.

Le but de cette phase est de caractériser les flux de chaque application (type, périodicité) et d'identifier les acteurs qui émettent et reçoivent les données.

Étape 1 — Identifier les flux

Application	De ↔ vers	Objet	Type de flux	Système	Périodicité
Base de données	Directions régionales ↔ siège	Comptabilité, logistique	Client-serveur	Unix	Mise à jour : TLN* Consultation : TLJ*
Datawarehouse	Directions régionales ↔ siège	Activité commerciale	Transfert de gros fichiers	Unix	Mise à jour : TLN Consultation : TLJ
Messagerie	Intrasite et intersite		Transfert de fichiers (Word, Excel, appli- cations métier)	Exchange SMTP	Toutes les 10 minutes entre MTA (un MTA par site)
Télécopie via la messagerie	Tous les sites	Échanges externes	Messagerie	Exchange	TLJ
Serveur web	Tous les sites	Informations, DRH, accès aux bases de données	Transactionnel Client-serveur	Unix, NT	TLJ
Connexions internet	Directions régionales ↔ Siège	Consultation web messagerie	Transactionnel, transfert de fichiers	Accès Internet situé au siège	TLJ

*TLJ = tous les jours

*TLN = toutes les nuits

Cette vision synthétique est une étape vers la traduction du langage utilisateur en langage informatique. C'est aussi un bon moyen de décrire les flux circulant au sein de la société (*workflow*) afin de bâtir le réseau qui lui soit le mieux adapté.

Les flux recensés peuvent être classés en trois catégories :

- les flux conversationnels ;
- les flux transactionnels ;
- les flux de type transfert de fichiers.

Il faut ajouter à cela les applications client-serveur qui peuvent, selon les cas, s'apparenter à la deuxième ou à la troisième catégorie.

Quel type de flux ?	Quelles caractéristiques ?	Quelles applications ?
Conversationnel	Trames courtes, fréquence soutenue	Connexions Telnet
Transactionnel	Trames moyennes (trafic montant) et longues par rafales (trafic descendant)	Serveurs web Connexions aux sites centraux (via des passerelles)
Transfert de fichiers	Trames longues, trafic soutenu	Serveurs bureautique (FTP ou moniteur spécialisé sur TCP/ IP)
Client-serveur	Dépend de la position de la base de données et du module client	Requêtes SQL sous Unix, Windows NT, etc.

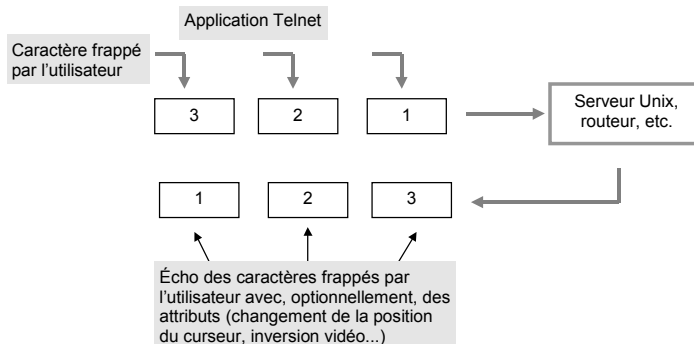
Ces flux doivent cohabiter au sein d'un même réseau intersite et être transportés simultanément sur une même liaison.

Les flux de type conversationnel

Les applications conversationnelles sont les plus courantes dans les mondes Unix et TCP/IP. Le protocole utilisé est *Telnet*. Le principe repose sur l'envoi d'un caractère avec écho distant. Une session étant établie entre un poste de travail et une machine, tout caractère frappé sur le clavier est envoyé à la machine, traité par cette dernière, et enfin renvoyé tel quel pour affichage, éventuellement avec d'autres attributs. Chaque caractère peut en effet déclencher une action comme l'affichage d'une fenêtre.

Figure 9-11.

*Types de flux
générés
par des applications
conversationnelles.*



Le type de flux qui en résulte est par conséquent irrégulier car il dépend de l'activité de l'utilisateur et est composé de trames courtes.

Le temps de réponse est donc primordial pour ce type d'application. Il se doit d'être le plus régulier possible, le principe étant qu'un utilisateur s'habitue à un temps de réponse, même mauvais, pourvu qu'il soit régulier. Un maximum de 300 à 500 ms est généralement toléré. Quand plusieurs caractères sont saisis à la suite, ce temps est généralement réduit du fait de leur encapsulation dans le même paquet TCP (algorithme de Nagle — RFC 896).

! Activation de l'algorithme nagle sur un routeur Cisco

! Utile pour les connexions Telnet

! A désactiver pour X-Windows

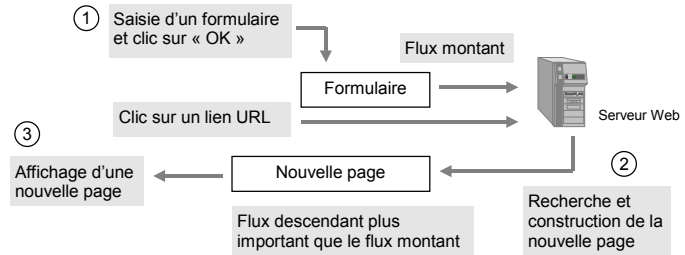
service nagle

Les flux de type transactionnel

Le mode transactionnel est le fonctionnement le plus courant pour les applications critiques sur les systèmes centraux. La technique consiste à envoyer un écran de saisie vers un terminal, à réaliser localement les modifications, puis à renvoyer les données modifiées vers le site central. Ces opérations sont contrôlées par un logiciel appelé moniteur transactionnel (CICS sous IBM et Tuxedo sous Unix).

Les flux générés entre serveurs web et navigateurs peuvent être assimilés au mode transactionnel, bien que le volume des pages web soit beaucoup plus important.

Figure 9-12.
*Types de flux générés
par des applications web.*

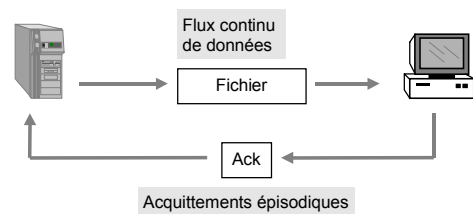


Les flux générés sont caractérisés par un trafic descendant (serveur web vers navigateur) plus important que le trafic montant (les données du formulaire ou un clic sur une URL). La ligne est rarement mobilisée (deux à quatre transactions par minute) tandis que le transfert d'une page (de 4 à 50 Ko, voire plus) requiert la presque totalité de la bande passante pendant quelques secondes. Le débit instantané requis est donc une donnée importante dans le calcul de la bande passante requise par ce type de flux.

Les flux de type transfert de fichiers

Ces flux sont caractérisés par des échanges soutenus et des trames longues. Leurs occurrences peuvent être prévisibles, dans la mesure où la majorité des transferts de fichiers est souvent associée à des traitements automatiques qui ont lieu en dehors des heures ouvrées, par exemple lors de la sauvegarde ou de la synchronisation de bases de données.

Figure 9-13.
*Types de flux générés
par des applications
transactionnelles.*



Pendant les heures ouvrées, ce type de flux peut dégrader les temps de réponse des flux transactionnels et surtout des flux conversationnels. Cette interférence peut être contrôlée par des mécanismes de priorité positionnés sur les équipements d'interconnexion tels que les routeurs.

Les flux client-serveur

Le concept client-serveur se décline en réalité en plusieurs modèles :

- La base de données et la logique applicative sont situées sur le serveur. Le poste client soumet une requête puis attend les résultats qui, seuls, transitent par le réseau.
- Le serveur héberge la base de données ; la logique applicative réside sur le poste client. La puissance de traitement est donc reportée sur les postes client. Les échanges sur le réseau sont aussi fréquents que les manipulations de la base.
- La logique applicative et les données sont réparties entre le serveur et le client. La procédure d'interrogation consiste à extraire tout ou partie de la base de données centrale, puis à opérer des traitements spécifiques sur la base de données locale. Le réseau n'est sollicité que lors des extractions depuis la base de données. La synchronisation des bases peut intervenir en dehors des heures ouvrées.

Modèle	Application	Flux réseau
1. Base de données et application sur le serveur	Comptabilité, gestion commerciale, logistique, gestion des ressources humaines	Type transactionnel
2. Base de données sur le serveur et applications sur le client	Gestion commerciale Serveur web	Type transactionnel (volumétrie faible) ou transfert de fichiers (volumétrie élevée)
3. Base de données et applications réparties entre le serveur et le client	Datawarehouse Architectures 3 tiers	Transfert de fichiers Transactionnel

L'architecture 3 tiers (client, serveur applicatif, serveur de base de données) génère des flux de type transactionnel côté client et de type transfert de fichiers côté base de données.

Estimer la volumétrie

Les flux doivent ensuite être quantifiés, soit à partir de données existantes, soit sur la base d'hypothèses. Si on part d'un réseau existant, soit pour l'optimiser, soit pour le faire évoluer, on peut s'appuyer sur des statistiques indiquant les volumes échangés entre deux sites. Ces données peuvent être issues de facturations détaillées ou d'une phase d'audit consistant en une campagne de mesure sur le terrain.

La volumétrie est calculée différemment selon le type de flux. Souvent, elle doit être extrapolée à partir d'informations partielles. Ce travail doit donc être réalisé indépendamment pour chaque application que le réseau intersite sera susceptible de véhiculer. Les résultats doivent ensuite être consolidés sous forme de matrice de flux présentant les volumes échangés entre chaque site. L'échelle de temps généralement utilisée est une journée de travail ; cette périodicité permet en effet de lisser les variations.

La volumétrie globale pour un site est généralement issue d'une volumétrie unitaire estimée pour un utilisateur et calculée selon la formule suivante :

$$V_j = V_u \times U$$

V_j est le volume journalier à calculer pour un site.

V_u est le volume journalier estimé pour un utilisateur.

U est le nombre d'utilisateurs pour un site donné.

Les sections suivantes décrivent les manières d'estimer les volumétries lorsque l'existant est peu ou pas connu.

Étape 2 — Estimer la volumétrie

Applications	Exemples d'estimation de la volumétrie
Messagerie SMTP ou Exchange	10 messages par utilisateur et par jour x 100 Ko Synchronisation des annuaires
Transfert de fichiers FTP	N % des utilisateurs (ou des applications batch) = X Ko par jour
Transactionnelles sites centraux	100 à 200 écrans de 2 ou 4 Ko par utilisateur et par jour
Transactionnelles web	20 à 50 écrans de 4 à 50 Ko par utilisateur et par jour
Conversationnelles Telnet	Dépend des applications (faire des tests avec un analyseur réseau) ; un écran = 2 à 4 Ko
Services réseaux (vidéotex, télécopie, etc.)	3 sessions vidéotex par jour et par utilisateur Un fax de 10 Ko par jour
Administration du réseau	0 à 10 sessions Telnet sur chaque routeur Configuration SNMP de 1 Ko par équipement et par jour 1 à 50 trap SNMP de 1Ko par jour Sondes RMON : 1 à 10 transferts de fichiers par jour (plusieurs centaines de Ko)

Volumétrie liée à la messagerie

Les volumes de données générés par une messagerie bureautique peuvent être modélisés sur la base des hypothèses suivantes :

- Environ dix messages par jour et par utilisateur à destination d'un autre site (20 % des messages sont à destination d'un site extérieur, 80 % restent locaux).
- Environ 100 Ko par message. Cette valeur dépend beaucoup de l'utilisation qui est faite de la messagerie au sein de la société. Plus celle-ci est utilisée, plus les messages ont tendance à être importants (pièces jointes).
- La taille de l'annuaire est basée sur 100 octets par utilisateur.
- Synchronisation hebdomadaire (voire toutes les nuits) de l'annuaire : transfert depuis les sites distants vers le siège (si la gestion est décentralisée), consolidation de l'annuaire, puis transfert depuis le siège vers les sites distants.

Les messageries bureautique transportent les messages sous forme de copies de fichiers entre les serveurs bureautique. La périodicité des échanges dépend du paramétrage ; elle est généralement comprise entre 5 et 15 minutes. Ces transferts de fichiers occupent donc régulièrement la bande passante des liens.

Volumétrie liée aux transferts de fichiers

La volumétrie liée aux transferts de fichiers dépend des applications présentes au sein de la société. Son évaluation repose donc sur une analyse précise de l'existant et/ou des besoins. Elle peut être modélisée sous la forme $N\%$ des utilisateurs réalisant l'équivalent d'un transfert de X Ko par jour à destination d'un site distant.

Volumétrie liée aux applications transactionnelles site central

Dans la plupart des cas, on peut estimer qu'un utilisateur échange 100 à 200 écrans de 2 Ko à 4 Ko par jour avec le site central. Cette évaluation est bien sûr éminemment variable selon le contexte à considérer. La taille des écrans varie, par exemple, en fonction des applications, et la fréquence des échanges en fonction du type de travail de l'utilisateur (saisie intensive, consultation, etc.). Il convient donc d'estimer la volumétrie moyenne à partir de tests.

Volumétrie liée aux applications transactionnelles web

Même remarque que pour les applications transactionnelles, sauf que la taille des pages varie entre 4 Ko et 50 Ko, une page pouvant contenir des images GIF (fixes ou animées).

En prenant en compte les fichiers GIF, JPG et HTML, la moyenne constatée est de 4 Ko. Si on prend en compte les transferts de fichiers réalisés à partir de l'Internet (documents .pdf, .txt ou .doc), la moyenne constatée est de 100 Ko. La moyenne peut atteindre plusieurs mégaoctets si le téléchargement des exécutables (.exe) est autorisé.



Vous pouvez vérifier les valeurs propres à votre contexte en visualisant le contenu du cache de votre navigateur (recherchez un répertoire appelé "cache" situé dans le répertoire d'installation du navigateur).

Volumétrie liée à d'autres services

Différents services peuvent emprunter le réseau intersite, notamment en provenance de sites rattachés dans le cas où les passerelles de communication sont centralisées. Les hypothèses de travail qui peuvent être retenues sont les suivantes (il ne s'agit ici que d'indications, la volumétrie réelle étant liée à la nature des travaux réalisés par les utilisateurs) :

- Service de télécopie. Chaque utilisateur expédie en moyenne un fax de 10 Ko par jour.
- Service d'accès au vidéotex. Un quart des utilisateurs effectuent trois connexions vidéotex de 2 minutes par jour. Chaque connexion génère un flux de 5 Ko.
- Etc.

L'utilisation des applications multimédias pose d'autres problèmes, qui sont abordés au chapitre 14.

Rassembler toutes les données

Ayant ces abaques en tête, nous pouvons maintenant calculer les volumes pour notre cas. La première chose à faire est d'établir la matrice des flux qui présente les types de flux et le nombre d'utilisateurs qui les génèrent.

Étape 3 — Matrice de flux

Depuis \ vers	Toulouse	Paris	Strasbourg	Etc.
Toulouse	---	400 web 40 Telnet	100 Telnet	
Paris		---		
Strasbourg	50 web		---	
Etc.				---

La volumétrie doit être calculée entre chaque site et dans les deux sens. Les liaisons étant de type *full duplex*, il convient de prendre la valeur la plus haute, ce qui permet de calculer le débit instantané nécessaire.

Étape 4 — Matrice volumétrique

	De → vers				
	P → T	T → P	T → S	S → T	etc.
Flux transactionnels web					
Nombre d'utilisateurs	400	400	50	50	...
Volumétrie unitaire (en Ko)	50	0,5	50	0,5	...
Pages par jour et par utilisateur	10	10	10	10	...
Quantité / utilisateur / jour (en Ko)	500	5	500	5	...
Total journalier en Mo	200	2	25	0,25	...
Flux conversationnels Telnet					
Nombre d'utilisateurs	40	40	100	100	...
Volumétrie unitaire (en Ko)	4	0,5	0,5	4	...
Écrans par jour et par utilisateur	100	100	50	50	...
Quantité / utilisateur / jour (en Ko)	400	50	25	200	...
Total journalier en Mo	16	2	2,5	20	...
Etc.	...	...	...	...	...
Volume total en Mo	960	320	240	190	

Selon le sens de la connexion client-serveur, les flux montants (depuis le client vers le serveur) et descendants (depuis le serveur vers le client) apparaîtront dans la première colonne ou la deuxième.

Au final, seul le maximum des deux flux doit être pris en compte. C'est lui qui déterminera la bande passante maximale requise.

Calculer les débits

Pour dimensionner une liaison, il convient d'estimer les besoins en termes de débit instantané. La formule de calcul généralement admise est la suivante :

$$Bp = Vj \times Th \times Ov \times \frac{1}{Tu} \times \frac{1}{3600} \times (8 \times 1,024)$$

La signification des paramètres est la suivante :

- Bp** est la bande passante instantanée calculée pour une liaison exprimée en Kbit/s.
- Vj** est le volume journalier, estimé en Ko. Cette valeur représente la somme des flux devant circuler sur le lien considéré (le maximum pris entre les flux montants et descendants)
- Th** est un coefficient permettant de calculer le trafic ramené à l'heure chargée. On considère généralement que le trafic journalier est concentré sur une heure chargée. Cette hypothèse part du constat que, sur 8 heures de travail, les utilisateurs sont le plus actifs sur deux périodes de pointe, entre 10 h et 11 h, et entre 15 h et 16 h. Les valeurs généralement admises sont comprises entre 20 % et 30 % du trafic journalier concentré sur une heure.
- Ov** est l'*overhead* généré par les protocoles de transport (TCP, IP, PPP). Ce coefficient est généralement affecté d'une valeur de 20 %. Il tient compte des en-têtes et des paquets de service (acquittements, etc.).
- Tu** est le taux maximal d'utilisation de la bande passante du lien. Cette correction permet de prendre en compte le fait que l'on utilise rarement 100 % du débit nominal d'un lien. Ce taux est généralement fixé à 80 % de la bande passante, ce qui donne un surdimensionnement du lien de l'ordre de 25 %. Pour des liaisons à haut débit, ce taux peut atteindre 90 %.

Le rapport 1/3600 permet de ramener la volumétrie sur une heure en secondes, tandis que le rapport 8*1,024 permet de convertir les kilo-octets en kilobits (1 octet = 8 bits, 1 Ko = 1 024 octets et 1 000 bits = 1 kilobit).

Si on prend les valeurs standard pour ces paramètres, la formule devient :

$$B_p = V_j \times 0,30 \times 1,2 \times \frac{1}{0,8} \times \frac{1}{3600} \times (8 \times 1024)$$

soit, par exemple, une bande passante de 1 Mbit/s pour un volume journalier estimé à 1 Go.

Si la liaison doit servir de secours pour n autres liaisons de débit D_N sans que les performances ne soient dégradées, la bande passante du lien doit être augmentée de la somme de ces débits D_N . Dans notre cas, nous choisirons un mode dégradé, afin de limiter les coûts.

Liaison	Volume en Ko V_j	30 % à l'heure chargée T_h	Overhead protocole O_v	Taux d'occupation du lien T_u	Débit du lien en Kbit/s
P ↔ T	960 000	0,3	1,2	0,8	984
T ↔ S	240 000	0,3	1,2	0,8	246
Etc.	...	...	...	...	...

Le débit du lien doit être arrondi à la valeur supérieure des débits proposés par les opérateurs, soit, dans notre cas, 1 Mbit/s entre Paris et Toulouse et 256 Kbit/s entre Toulouse et Strasbourg.

Tenir compte des temps de réponse

Pour des applications client-serveur reposant sur des extractions de données et assimilables à des transferts de fichiers, le critère performance se pose en termes de délai maximal de transfert des données. Il convient donc de calculer les débits nécessaires en fonction des délais acceptables et des volumes estimés :

$$B_p = \frac{V_o}{T_{ps}}$$

B_p est la bande passante nécessaire.

V_o est le volume moyen (converti en kilobits) des données extraites suite à une requête.

T_{ps} est le temps de réponse souhaité.

Cette démarche est à combiner avec une étude de coût, car il faut trouver un compromis avec la performance. Il faut donc recourir à une simulation des temps de réponse obtenus en fonction des débits des liens, et éventuellement les mesurer pour différentes tailles de requêtes.

Par exemple, le tableau suivant compare les temps de transfert de données de différentes tailles en fonction du débit de la ligne.

Volume en Ko	64 Kbit/s	128 Kbit/s
10	1,25 s	0,62 s
20	2,50 s	1,25 s
30	3,75 s	1,87 s
50	6,25 s	3,12 s
Coût mensuel HT	7 844 F	14 264 F

Il vous appartient alors de mettre en balance le coût et les performances souhaitées.

10

Bâtir un réseau de transport

Le réseau WAN qui interconnecte les réseaux LAN utilise les services d'un **réseau de transport**. Celui-ci véhicule également de la voix, de la vidéo, etc.

Le réseau de transport correspond aux couches physiques (niveau 1) et logiques (niveau 2). Sur le LAN nous avons Ethernet, sur le WAN nous allons avoir ATM et Frame Relay.

Ces technologies répondent à des contraintes plus larges que celles d'une interconnexion de LAN. Notre réseau intersite n'est donc qu'un utilisateur parmi d'autres, tels qu'un réseau de PABX pour la téléphonie ou des connexions entre salles de visioconférence.

Dans ce chapitre, vous apprendrez ainsi :

- à interconnecter des réseaux locaux *via* Frame Relay et ATM ;
- à configurer les circuits virtuels ;
- à gérer la qualité de service ;
- à connaître la signalisation et l'adressage.

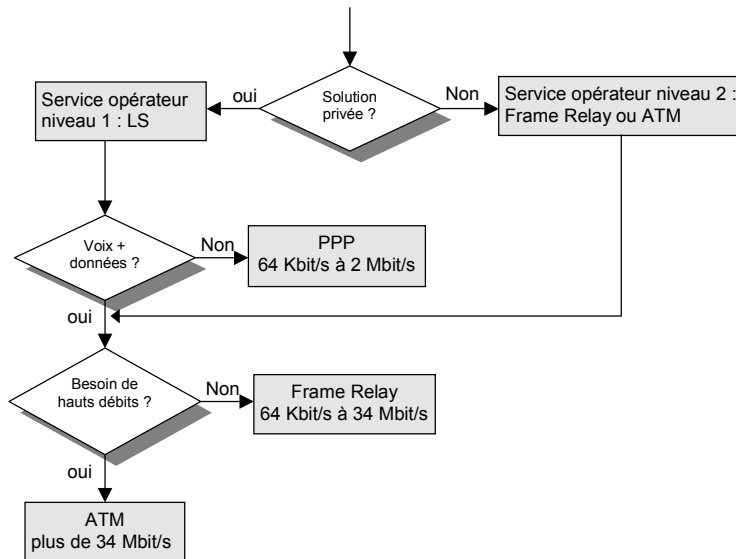
LS, Frame Relay ou ATM ?

En tant qu'utilisateur, seul le service compte, quelle que soit la technologie employée. Le choix de l'une des trois solutions qui s'offrent à nous se fera en fonction de nos besoins, de l'offre du marché et du coût, que ce soit dans le cadre d'une solution privée ou opérateur.

Comme nous l'avons vu au chapitre précédent, il est possible de choisir entre différentes technologies et différents niveaux de service.

Pour une interconnexion de réseaux locaux, la LS est la solution idéale. À partir d'un certain nombre de sites (une dizaine, voire moins à l'international), la solution opérateur reposant sur un réseau Frame Relay ou ATM devient plus rentable.

La technologie Frame Relay est actuellement la plus répandue, mais les opérateurs investissent dans ATM. En 1997, leurs réseaux reposaient pour 30 % sur une infrastructure Frame Relay et pour 45 % sur ATM. En 1998, la proportion était de 22 % pour Frame Relay et 61 % pour ATM.



Pour l'utilisateur, un accès Frame Relay revient cependant moins cher. Les opérateurs réservent donc ATM pour des accès à haut débit (34 Mbit/s et plus), et limitent les accès Frame-Relay à 34 Mbit/s (plus rarement à 45 Mbit/s).

Techniquement, ATM a tout pour s'imposer, mais l'histoire de l'informatique nous a montré que cela n'est pas un gage de pérennité. Rappelons-le, les protocoles qui se sont imposés sont les plus simples, les moins chers et surtout les mieux adaptés aux besoins des utilisateurs.

Par exemple, les principaux inconvénients d'ATM sont :

- un *overhead* très important (plus de 10 %) ;
- un manque de maturité dans la normalisation et les offres.

En revanche, le point fort d'ATM réside dans sa capacité multiservice : il permet de créer des réseaux locaux et étendus (LAN et WAN) et transporte les flux multimédias. Malheureusement, c'est justement sur ces points qu'ATM manque de maturité : les normes ne sont pas stabilisées, ou non satisfaisantes, et aucune offre ne permet de transporter simultanément les flux voix, données et vidéo.

En réalité, ATM est actuellement utilisé en tant que :

- réseau fédérateur WAN pour les réseaux fédérateurs (backbone) des opérateurs ;
- réseau fédérateur LAN par les entreprises, en association avec les réseaux virtuels LANE (*LAN Emulation*).

Mais, même sur ces créneaux, il est concurrencé par le Gigabit Ethernet.

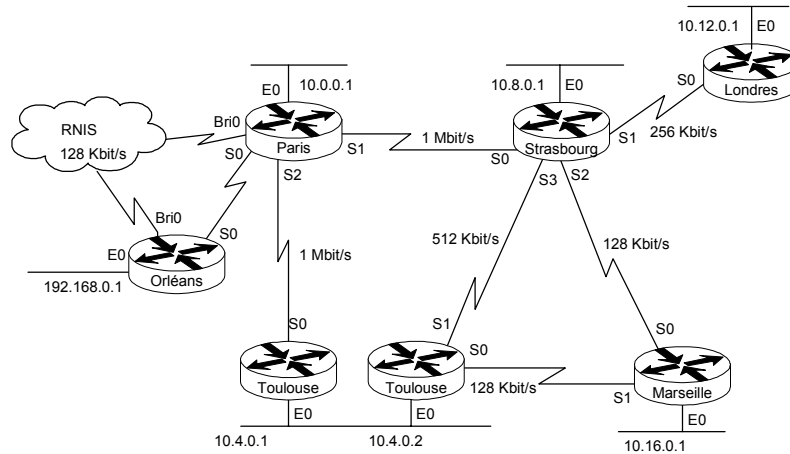
Critère	Frame Relay	ATM
Débit	De 64 Kbit/s à 45 Mbit/s	À partir de 34 Mbit/s
Qualité de service	Gestion des congestions, débit garanti	Gestion des congestions, débit garanti, trafic synchrone, priorités
Réseau	WAN	LAN et WAN
Application	Voix et données (vidéo via IP)	Voix, données et vidéo
DSU	FRAD	DXI ou ATM
Overhead pour un MTU de 1 500 octets	0,5 %	10,4 % au minimum
Adressage	Local (DLCI)	Local (VPI/VCi)
Normes	ITU Q.922 / Q.933	ITU I.361 à I.363 / Q.931 et ATM Forum
Transport d'IP	RFC 2427 (NLPID/SNAP)	RFC 1483 (LLC/SNAP) RFC 2225 (Classical IP)

À côté de ces deux protocoles, nous retrouvons nos LS (lignes spécialisées) qui sont dans tous les cas la base d'un réseau de transport Frame Relay ou ATM. Nous pouvons même décider de bâtir notre réseau uniquement sur des LS, comme nous l'avons fait pour notre première interconnexion.

Mettre en place un réseau de LS

La mise en place d'un réseau de LS implique de commander nous même les liaisons auprès des opérateurs locaux. En France, on peut s'adresser à France Télécom et, bientôt, à d'autres pour toutes nos LS. Il suffit d'indiquer à l'opérateur le débit souhaité, les adresses des sites à connecter et, si cela est proposé, le type d'interface désirée : V.35 ou X21/V11 la plupart du temps.

Figure 10-1.
Réseau intersite
reposant sur des LS.



Pour la liaison internationale Strasbourg-Londres, le problème se complique puisque vous avez affaire à deux opérateurs, BT et France Télécom, par exemple. La liaison est administrativement découpée en **demi-circuits**, chacun étant géré par l'opérateur situé aux tenants et aboutissants de la LS.

Votre correspondant à Londres doit donc s'adresser à BT pour la fourniture du demi-circuit anglais, et vous devez faire de même auprès de France Télécom pour le demi-circuit français. Dans le cas où la LS traverse plusieurs autres pays, aucune autre démarche n'est à effectuer.

Les opérateurs ont passé des accords de coopération multilatéraux et ont mis en place des structures de coordination pour que les demi-circuits soient regroupés en une seule liaison internationale. Cette procédure est transparente pour vous. Néanmoins, le délai de mise en service est généralement de dix semaines, voire plus dans certains pays.

La mise en service de la LS se concrétise par l'installation d'un modem numérique CSU (*Channel Service Unit*) dans vos locaux et par un test BERT (*Bit Error Tests*). Ce dernier consiste à envoyer un train continu de bits pendant une durée minimale de 24 heures, et à mesurer le taux d'erreur qui ne doit pas dépasser 10^{-6} à 10^{-10} (une erreur tous les dix milliards de bits).

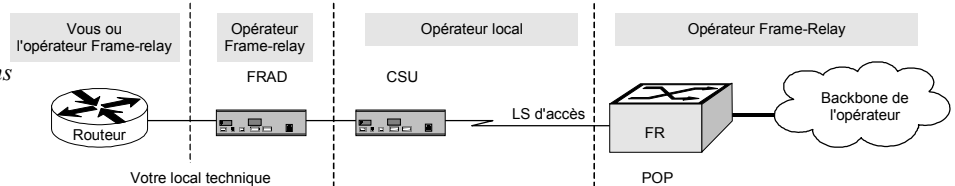
Généralement, l'opérateur vous accorde un délai de quelques jours avant de déclarer la liaison opérationnelle et de débiter la facturation. Cela vous permet de tester la liaison avec vos routeurs.

Mettre en place un réseau Frame Relay

Faire appel à un opérateur permet de ne s'adresser qu'à un seul interlocuteur, quels que soient les pays concernés. Même en retenant un opérateur étranger pour vos sites français, celui-ci se chargera de commander les liaisons d'accès (en fait, des LS aux technologies E1 ou xDSL) auprès de l'opérateur local. Il suffit de lui indiquer les adresses de vos sites pour qu'il commande les LS entre vos sites et ses POP les plus proches.

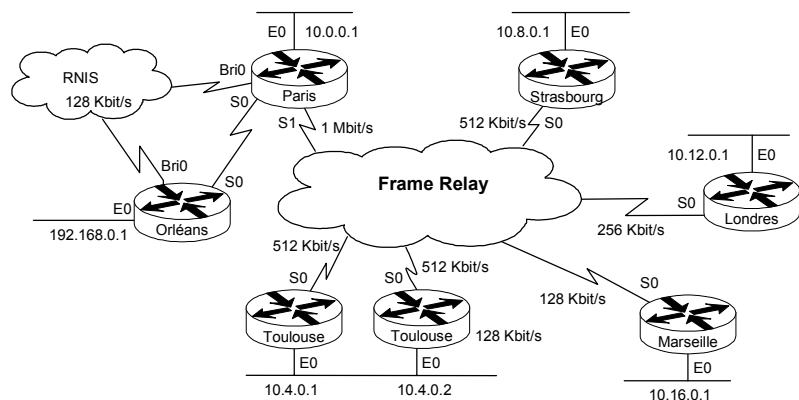
Comme pour la LS, l'opérateur local installe son modem numérique (le CSU) dans vos locaux. L'opérateur retenu pour le réseau Frame-Relay installe ensuite un commutateur appelé **FRAD** (*Frame-Relay Access Device*) qu'il va connecter au modem. L'interface série du routeur sera ensuite connectée à un des ports du FRAD.

Figure 10-2.
*Responsabilités
des configurations
de l'accès
Frame Relay.*



Le FRAD est un équipement de conversion entre des protocoles d'entrée (interfaces E1, X.21/V11, voix, données, etc.) et le protocole Frame Relay. Cet équipement prend les paquets IP issus d'un routeur, ou les canaux voix issus d'un PABX, et les encapsule dans les trames Frame Relay. Le FRAD est relié via la liaison d'accès au commutateur Frame Relay situé dans le POP de l'opérateur. Plusieurs flux sont ainsi multiplexés sur une même liaison.

Figure 10-3.
*Réseau intersite
reposant
sur un réseau opérateur
Frame Relay.*



Les avantages sur la première solution sont immédiatement perceptibles :

- Il n'y a plus qu'une seule LS par site. Nous avons besoin de moins d'interfaces série, ce qui diminue d'autant le coût de nos routeurs.
- Les LS sont locales entre nos sites et les POP de l'opérateur, ce qui diminue également leur coût.

Critère	Solution opérateur	Solution privée
Caractéristiques de l'offre	Un réseau fédérateur et des liaisons locales.	Des liaisons longue distance entre les sites.
Déménagement d'un site	Facile : une liaison locale à changer. Coûts réduits.	Difficile : plusieurs liaisons longue distance à changer. Coûts élevés dus aux changements et à la période de recouvrement.
Modification des débits	Très souple : débit des liaisons locales et/ou des CIR.	Assez difficile (surtout à l'international) : changement du débit des liaisons longue distance.
Redondance de site (second site en secours du premier)	Possible : une liaison locale sur le site principal et une sur le site de secours.	Impossible, sauf à doubler les liaisons longue distance (très cher).
Exploitation du réseau	Prise en charge par l'opérateur.	Prise en charge par le client <i>via</i> des contrats de maintenance.

Nous avons conservé notre LS entre Paris et Orléans, car le réseau opérateur était plus cher dans ce cas précis. Cela permet de montrer l'usage de trois réseaux d'opérateurs : liaison spécialisée, RNIS et Frame-Relay. Par ailleurs, deux liaisons d'accès ont été conservées à Toulouse car le site est stratégique.

Qualité de service et facturation

Un opérateur s'engage toujours sur une qualité de service et base son offre commerciale sur les mécanismes offerts par les protocoles qu'il utilise.

Par exemple, Frame Relay permet de **gérer la congestion** du réseau, de diminuer le débit des LS et donc les coûts. Comment ? La réponse est intimement liée au fonctionnement du protocole Frame Relay.

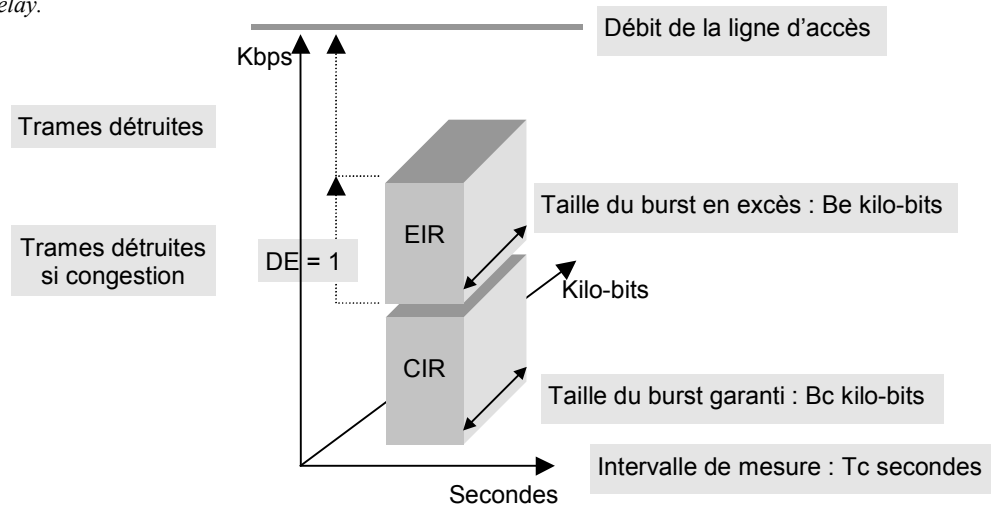
Débit garanti

Tout d'abord, le réseau garantit à l'utilisateur un volume de B_c (*committed burst size*) kilobits pendant une période de T_c (*committed rate measurement interval*) secondes, définissant ainsi un débit garanti **CIR** (*Committed Information Rate*) : $CIR = B_c / T_c$. Lorsque le commutateur voit passer plus de B_c kilobits pendant T_c secondes (c'est-à-dire lorsque le débit dépasse le CIR), le bit **DE** (*Discard Eligibility*) des trames en dépassement est positionné à "1". Cela signifie que ces trames seront détruites en priorité en cas de congestion du réseau.

Ensuite, le réseau autorise à l'utilisateur un volume supplémentaire de B_e (*excess burst size*) kilobits pendant une période de T_c secondes, définissant ainsi un débit en excédent **EIR** (*Excess Information Rate*) : $EIR = B_e / T_c$. Lorsque le commutateur voit passer plus de $B_c + B_e$ kilobits pendant T_c secondes (lorsque le débit dépasse l'EIR), toutes les trames en excès sont détruites.

Généralement, la période de mesure T_c est fixée à une seconde, et l'**AIR** (*Allowed Information Rate* = CIR + EIR) ne dépasse pas le débit de la liaison d'accès.

Figure 10-4.
Qualité de service
Frame Relay.



Le débit des liaisons internes au réseau peut être inférieur à la somme des AIR des clients. L'opérateur se fonde sur des calculs de probabilité qui montrent que tous les clients n'utiliseront pas leur AIR en même temps (*surbooking*), et compte sur les mécanismes de contrôle de congestion offerts par le protocole pour résoudre les problèmes. Cela permet à l'opérateur de ne facturer que le CIR (le débit garanti), l'EIR étant gratuit ou presque.

Dès lors, tout le monde "joue" sur les CIR : le client pour payer le moins cher possible, et l'opérateur pour dépenser le moins possible.

Le débit de la liaison d'accès peut ainsi être de 512 Kbit/s, et celui du CIR de 64 Kbit/s ; cela permet d'obtenir un débit maximal de 512 Kbit/s tout en ne payant que pour 64 Kbit/s. Inversement, pour un site central, la somme des CIR peut être égale à 200 % du débit de la liaison d'accès (on espère alors que toutes les applications n'utiliseront pas le réseau en même temps).

Cependant, le débit offert au-delà du CIR n'étant pas garanti, les temps de réponse risquent d'être mauvais et erratiques. Il est donc conseillé de dimensionner suffisamment les CIR afin d'obtenir une bonne qualité de service, notamment pour les applications que vous considérez comme étant critiques et surtout pour les flux voix.

Connecter un routeur au réseau de transport

La norme Frame Relay définit uniquement une **interface d'accès** pour l'utilisateur. Bien qu'il soit possible de réaliser un réseau 100 % Frame-Relay, le réseau interne de l'opérateur repose souvent sur un protocole propriétaire ou ATM. Pour le client, cela n'a pas d'importance : il faut simplement que l'interface soit de type Frame Relay.

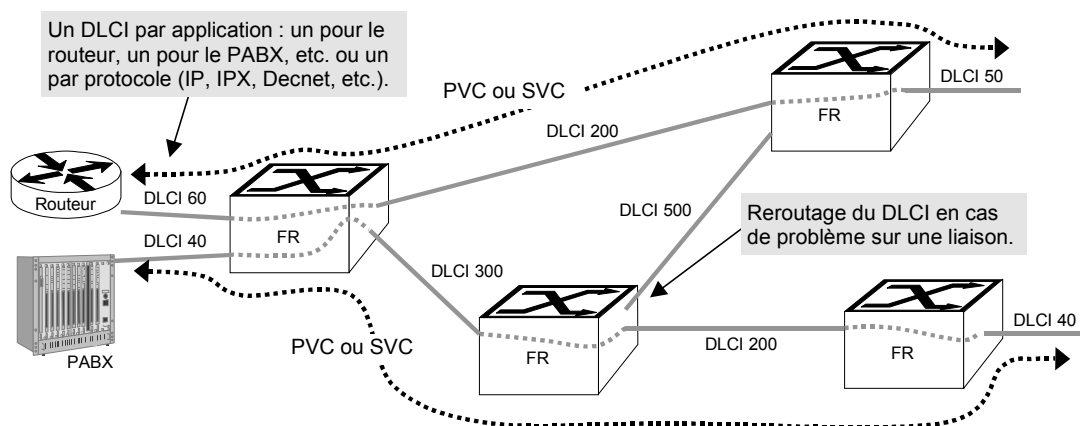
Chaque trame est identifiée par un **DLCI** (*Data Link Connection Identifier*), une adresse locale partagée par deux commutateurs. Il n'y a pas d'adressage de bout en bout, mais uniquement un adressage point à point entre deux commutateurs. Un commutateur recevant une trame avec un DLCI donné la routera sur un autre port et l'enverra avec un autre DLCI, et ainsi de suite.

La connexion entre deux commutateurs s'effectue par l'ouverture de circuits virtuels permanents (**PVC**, *Permanent Virtual Circuit*) ou commutés (**SVC**, *Switched Virtual Circuit*), c'est-à-dire ouverts à la demande *via* un protocole de signalisation (qui utilise quelques Kbit/s dans le DLCI 0).

L'intérêt des SVC est que la qualité de service (CIR et EIR) peut être spécifiée à la demande permettant ainsi de réduire (encore) les coûts. Sur un PVC, le CIR est fixé une fois pour toutes et engendre un coût fixe. De plus, la signalisation Q.933 utilisée pour établir les SVC permet de demander un délai de transit maximal, alors que cette fonction n'est pas prévue pour les PVC.

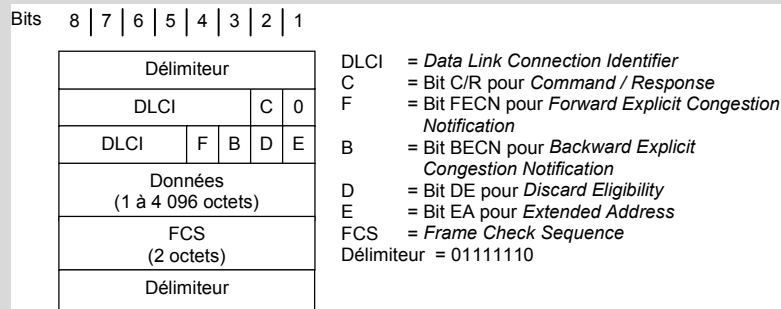
Figure 10-5.

*Circuits virtuels et DLCI
Frame Relay.*



LE POINT SUR FRAME RELAY (ITU Q.922 ANNEXE A, ANSI T1.618, FRF 1.1)

Le Frame Relay (relais de trames) est un protocole de niveau 2 multipoint qui définit uniquement l'interface d'accès au réseau (**UNI**, *User Network Interface*). Les nœuds intermédiaires relaient les trames sans réaliser le moindre contrôle de flux ni aucune reprise sur erreur : les trames peuvent ainsi être transportées par n'importe quel protocole. Seuls les nœuds d'extrémité sont tenus de respecter la norme Q.922, appelée "*Frame Relaying Bearer Service*".



Le champ FCS est un code de contrôle d'erreur de type CRC (*Cyclic Redundancy Code*). Si une erreur est détectée, la trame est détruite. Les pertes de données et les reprises sur erreur sont laissées à l'initiative des couches supérieures (TCP dans les cas Internet/intranet). Le bit EA permet d'étendre les 10 bits du DLCI à 16 ou 23 bits.

Lorsque le débit des trames atteint le **CIR** (*Committed Information Rate*), le commutateur positionne le bit **DE** à "1". Si, dans le réseau une **congestion** est décelée, les trames marquées DE seront détruites en priorité. Autrement dit, le débit peut donc dépasser le CIR, mais sans garantie.

Les congestions sont détectées au niveau des files d'attente. Lorsqu'un seuil est dépassé, le commutateur avertit explicitement l'émetteur du flux (en positionnant à "1" le bit **BECN** des trames circulant dans l'autre sens) et le récepteur du flux (en positionnant le bit **FECN** à "1").

La recommandation Q.922 suggère que le commutateur émetteur réduise son flux de 30 % si, pour plus de S % des trames qu'il reçoit, le bit BECN est positionné à "1". La valeur S est calculée dynamiquement en fonction du débit, du délai de transit, des paramètres Bc et Be, etc. Si le bit BECN est toujours positionné à "1" dans les trames qui continuent d'arriver, le flux est réduit de 50 %, puis de 75 % si le problème persiste.

La recommandation Q.922 suggère également que le commutateur récepteur réduise son flux de 25 % si, pour plus de la moitié des trames qu'il reçoit, le bit FECN est positionné à "1". Si la proportion s'inverse, il peut augmenter son flux par paliers de 1/16.

Les équipements terminaux (les routeurs, par exemple) peuvent également interpréter le bit BECN, et réagir de la même façon que les commutateurs. De même, ils peuvent détecter implicitement des problèmes de congestion lorsque des trames sont perdues après qu'un certain nombre d'entre elles aient été reçues avec le bit DE à "1".

Cependant, la norme Q.922 préconise que les commutateurs avertissent les équipements terminaux des congestions via le protocole **CLLM** (*Consolidated Link Layer Management*). Les messages CLLM (envoyés dans le DLCI 1023) contiennent la liste des DLCI pouvant causer des congestions à court, moyen et long termes. Cette signalisation de niveau 3 est plus fiable que l'interprétation des bits BECN et DE, car elle évite l'attente de trames (seuls vecteurs des bits BECN et DE) et évite aux couches supérieures d'avoir affaire à des informations internes au niveau 2.

Un PVC peut être géré comme un SVC au sein du réseau de l'opérateur. Cette facilité lui permet d'offrir le reroutage du PVC en cas de panne. Le client ne voit cependant qu'un PVC, appelé *soft PVC*, *switched PVC* ou encore *shadowed PVC*.

Si le routeur ne supporte pas Frame Relay

La manière la plus simple de raccorder nos routeurs est de les connecter à des FRAD (*Frame Relay Access Device*) de l'opérateur *via* leur interface série.

Une liaison Frame Relay peut être configurée en point à point ou en multipoint. Dans tous les cas, elle permet au routeur de joindre directement n'importe quel autre routeur. Elle doit donc être considérée comme étant un réseau IP, au même titre qu'un segment Ethernet, ce qui implique d'affecter une adresse IP à l'interface série du routeur.

```
interface serial 1
ip address 172.16.0.1 255.255.255.252
encapsulation hdlc
```

QU'EST CE QUE LA SIGNALISATION ?

Dans le jargon de l'ITU, la signalisation est un protocole qui permet de gérer l'état des connexions : ouverture et fermeture, négociation des paramètres, gestion de la qualité de service, etc. Vis-à-vis de l'utilisateur, la signalisation est appelée **UNI** (*User Network Interface*). Entre les équipements réseau, la signalisation est désignée sous le nom générique de **NNI** (*Network to Network Interface*).

Dans le monde TCP/IP, la séparation signalisation/données utilisateur n'est pas si nette : le trafic de service est souvent mêlé au trafic des données. Quelques protocoles utilisent cependant le principe de séparation, tel que FTP : le client et le serveur ouvrent deux ports TCP, le 21 pour envoyer les commandes, le 20 pour transférer les données.

Dans tous les cas, il s'agit d'une séparation logique, car la bande passante est partagée entre les données de service et les données utilisateur. Dans les réseaux des opérateurs, la signalisation peut cependant emprunter un chemin différent de celui des données.

En fait, cette séparation n'a réellement de sens que pour les protocoles de niveau 2 agissant en mode connecté, tels que Frame Relay et ATM. Une signalisation de niveau 3 permet ainsi de gérer les circuits virtuels de bout en bout à l'aide d'un adressage global de niveau 3.

Norme ITU	Description
Q.2931 et suivants	UNI ATM pour les réseaux publics = DSS2 (<i>Digital Subscriber Signalling System No. 2</i>)
Q.2100 et suivants	UNI ATM pour les réseaux privés = SAAL (<i>Signaling ATM Adaptation Layer</i>)
Q.931	UNI RNIS = DSS1 (<i>Digital Subscriber Signalling System No. 1</i>)
Q.933	UNI Frame Relay = DSS1 (<i>Digital Subscriber Signalling System No. 1</i>)

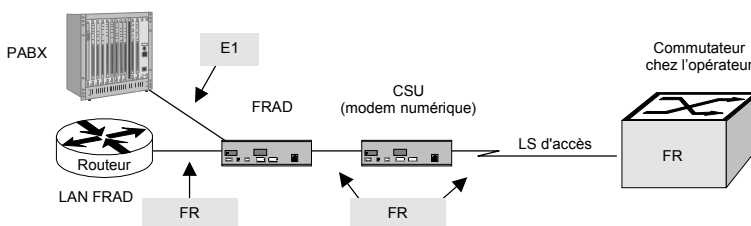
Toutes ces normes reposent sur le même protocole (format des messages, procédures, paramètres généraux). Seuls changent les paramètres propres à chaque réseau.

Si le routeur supporte Frame Relay

La solution la plus souple consiste cependant à configurer le routeur en FRAD : les trames LAN (Ethernet dans notre cas) sont converties en trames Frame Relay.

Dans ce cas, le FRAD de l'opérateur n'est utile que si d'autres équipements sont connectés, comme un PABX.

Figure 10-6.
*Connexion
d'un routeur
à un FRAD.*



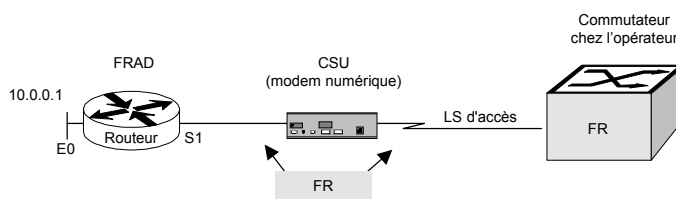
Le FRAD de l'opérateur se comporte comme un commutateur Frame-Relay vis-à-vis du routeur, et comme un FRAD voix (*voice-FRAD*) vis-à-vis du PABX.

Si vous choisissez un service de niveau 3, l'opérateur prend en charge le routeur qui peut alors être intégré au FRAD (ou inversement, le routeur peut supporter des cartes voix, permettant de configurer en *voice-FRAD*). Les constructeurs télécoms proposent des FRAD intégrant cartes voix et cartes routeur, tandis que les constructeurs informatiques proposent des routeurs intégrant des cartes voix.

Si, comme dans notre cas, seuls des réseaux locaux doivent être connectés, le routeur peut directement être raccordé au commutateur situé chez l'opérateur (dans son POP) :

```
interface serial 1
ip address 172.16.0.1 255.255.255.252
encapsulation frame-relay ietf
```

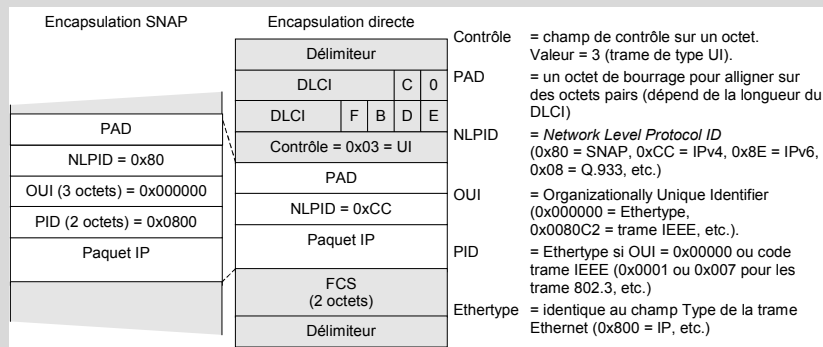
Figure 10-7.
*Connexion
d'un routeur
à un commutateur
Frame Relay.*



L'option "ietf" indique au routeur de respecter le RFC 2427 au lieu d'utiliser le format propre à Cisco.

ENCAPSULATION DANS FRAME-RELAY (RFC 2427, ANSI T1.617A, FRF 3.1)

Deux méthodes sont définies pour encapsuler les protocoles dans une trame Frame Relay. La plupart le sont par la méthode **SNAP** (*Sub Network Access Protocol*) tandis que certains, dont IP, peuvent l'être soit via SNAP, soit directement dans la trame via l'identificateur **NLPID** (*Network Level Protocol ID*).



Si la fonction pont est activée sur le routeur, les trames Ethernet, Token-Ring, FDDI, etc. sont encapsulées dans une trame Frame Relay selon la méthode SNAP.

Gérer les circuits virtuels

La gestion des CV (circuits virtuels) implique l'activation de trois mécanismes :

1. un protocole de signalisation Q.933 qui permet d'ouvrir et de fermer les SVC et de gérer les CV (PVC et SVC) ;
2. l'utilisation de la procédure LAP-F (décrite dans la norme Q.922) qui permet de transporter les messages Q.933 de manière sûre (reprise sur erreur, contrôle de flux, etc.) ;
3. un adressage global de bout en bout E.164 ou X.121 pour les SVC.

Ainsi, les routeurs peuvent échanger des informations avec les commutateurs en utilisant la signalisation LMI (*Local Management Interface*). Notre routeur supporte trois formats de messages : celui de l'ANSI (T1.617), celui de l'ITU (Q.933) et un autre, propre à Cisco, qui utilise le DLCI 1023 (en principe réservé à CLLM). Nous préférons utiliser le mode de fonctionnement défini par l'ITU.

```
interface serial1
frame-relay lmi-type q933a
```

"a" comme Annexe A de la norme Q.933

Les paramètres par défaut peuvent être utilisés pour la signalisation Q.933 et la procédure LAP-F. Cela simplifie la configuration. Si la liaison est de mauvaise qualité (souvent suite à un problème survenu sur la liaison d'accès), il peut être intéressant de réduire la taille initiale de la fenêtre à 8, voire à 1 :

```
interface serial1
frame-relay lapf k 8
```

LA SIGNALISATION FRAME RELAY (ITU Q.933, ANSI T1.617)

Les commutateurs Frame Relay utilisent un PVC dédié (DLCI 0) pour véhiculer les messages de signalisation Q.933. Ce protocole s'appuie sur celui utilisé par le RNIS (Q.931).

La première fonction offerte permet à l'équipement d'extrémité (un routeur, par exemple) de demander l'ouverture et la fermeture des **SVC** (*Switched Virtual Circuit*). Pour l'ouverture, les messages Setup et Connect contiennent le numéro d'appel (adresse aux formats **E.164** ou **X.121**), le DLCI affecté, etc., ainsi que le délai maximal de transit négocié. L'ouverture du SVC se traduit par l'affectation des DLCI entre les commutateurs et les équipements terminaux.

La deuxième fonction a trait à la procédure **LMI** (*Local Management Interface*) dont le rôle est de surveiller l'état des **PVC** (*Permanent Virtual Circuit*). Des messages *Status Enquiry* sont périodiquement envoyés par l'équipement terminal (un routeur, par exemple) pour connaître l'état du PVC. Le commutateur répond par un message *Status* (PVC actif, inactif, etc., ou nouveau PVC). En cas de défaillance, le routeur peut ainsi rerouter les flux. Le commutateur peut également interroger l'équipement terminal.

Les messages **ELMI** (*Extended LMI*) permettent au commutateur de communiquer à l'équipement terminal (un routeur, par exemple) la valeur des CIR, Be, Bc, etc. Le support de cette norme évite d'avoir à configurer les paramètres en double, à la fois sur le commutateur et sur le routeur.

Combien de circuits virtuels ?

La première question à se poser est de savoir combien de PVC ou de SVC utiliser : un circuit virtuel (CV) par protocole, ou un seul circuit virtuel pour tous les protocoles ?

- Un CV par protocole permet de bénéficier d'un débit garanti, ce qui assure que les transferts de fichiers ne perturberont pas les flux conversationnels. Mais cette solution coûte plus cher, car les opérateurs facturent chaque CV en fonction du CIR affecté.
- Un seul CV est plus économique, mais tous les flux de réseaux locaux sont mêlés. Le contrôle du flux doit donc être reporté au niveau du routeur à travers un système d'affectation de priorité par protocole ou de réservation de ressources (voir chapitre 14). Ces mécanismes fonctionnent généralement bien.

Nous choisirons donc la seconde solution.

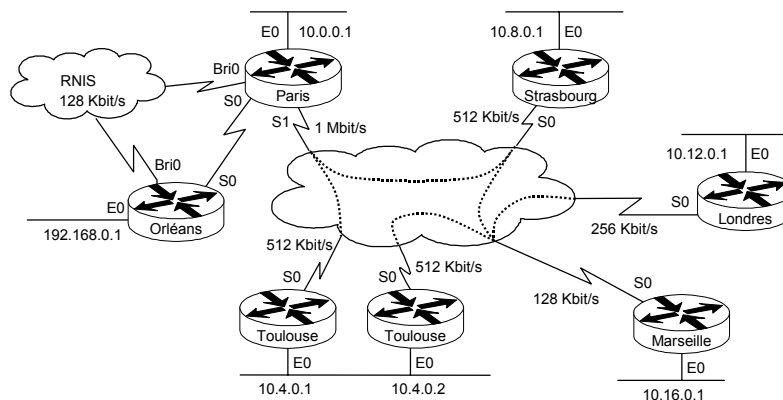
Un site peut communiquer avec plusieurs autres sites, ce qui implique d'utiliser un CV par connexion, les CV Frame Relay étant de type point à point.

Si nous voulons construire un réseau parfaitement maillé, chaque routeur doit voir un CV par site distant. Cette configuration est envisageable, mais risque d'être onéreuse, car l'opérateur facture chaque CV.

Si vous voulez réduire les coûts, il vaut mieux définir des CV là où les flux sont les plus importants. Les autres communications transiteront éventuellement par plusieurs routeurs (par exemple, Marseille ↔ Paris dans notre cas).

Figure 10-8.

Exemples
de circuits virtuels.



La question du choix entre PVC et SVC ayant été discutée précédemment, nous commencerons par configurer des PVC.

LA COMMUTATION FRAME RELAY (ITU Q.922)

La norme définit un second mode de fonctionnement, appelé “*Frame Switching Bearer Services*”, qui offre des mécanismes de reprise sur erreur et de contrôle de flux. Il repose sur l'utilisation de la procédure **LAP-F** (*Link Access Procedure-Frame*) inspirée du LAP-D du RNIS (Q.921), lui-même issu du LAP-B de HDLC. Le format des trames diffère légèrement, mais les principes restent les mêmes.

Ce mode de fonctionnement n'est actuellement utilisé que pour le PVC véhiculant la signalisation **Q.933** (dans le DLCI 0).

Bits	8	7	6	5	4	3	2	1
------	---	---	---	---	---	---	---	---

Octets	1	Délimiteur							
	2	DLCI						C	0
		DLCI		F	B	D	E		
	Contrôle (1 ou 2 octets)								
	Données (1 à 2 048 octets)								
	FCS (2 octets)								
Délimiteur									

DLCI	=	Data Link Connection Identifier
C	=	Bit C/R pour Command / Response
F	=	Bit FECN pour Forward Explicit Congestion Notification
B	=	Bit BECN pour Backward Explicit Congestion Notification
D	=	Bit DE pour Discard Eligibility
E	=	Bit EA pour Extended Address
FCS	=	Frame Check Sequence
Délimiteur	=	01111110

...

LA COMMUTATION FRAME RELAY (SUITE)

Le champ "Contrôle" a été ajouté à la trame *standard*. Il permet à la procédure LAPF de fournir un mécanisme de **reprise sur erreur** et de **contrôle de flux**. Il existe trois formats de trames (I, S, U) définissant autant de formats du champ de contrôle.

Bits

8

7

6

5

4

3

2

1

Trame I

N(S)

0

N(R)

P/F

Trame S

0

0

0

0

Su

Su

0

1

N(R)

P/F

Trame U

M

M

M

P/F

M

M

1

1

Trame I = Information

Trame S = Supervision

Trame U = Unnumbered

N(S) = Numéro séquence émission

N(R) = Numéro séquence réception

P/F = Poll/final (trame commande/réponse)

Su = Supervision

M = Modifier fonction

Les bits Su et M distinguent les 11 types de trames (SABME, UI, DISC, XID, RR, RNR, etc.)

Format de trame	Type	Description
I	--	Transport des données utilisateur
S	RR	<i>Receive Ready</i> : accusé de réception, et prêt à recevoir
	RNR	<i>Receive Not Ready</i> : accusé de réception, et non prêt à recevoir
	REJ	Reject : trame rejetée ; retransmettre à partir de N(R)
U	SABME	<i>Set Asynchronous Balanced Mode Extended</i> : initialisation de la procédure
	FRMR	<i>Frame Reject</i> : trame rejetée, erreur non récupérable
Etc.		

Afin de réduire le trafic de service, une **fenêtre d'émission** d'une taille *N* (entre 1 et 127) indique que le *commutateur* enverra *N* trames d'affilée et attendra un acquittement en retour. La taille de la fenêtre varie au cours du temps selon la qualité de la transmission (meilleure elle est, plus *N* est grand).

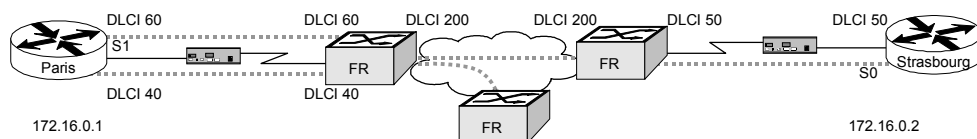
Le champ *N_A(S)* est le **numéro de séquence** de la trame envoyée par le commutateur A ; *N_A(R)* est celui de la trame *dernièrement* reçue du commutateur B, augmenté de 1. Si, après *N* trames envoyées, le commutateur A reçoit un *N_B(R)* inférieur à son *N_A(S)*, cela signifie que des trames ont été perdues ou endommagées. Il retransmet alors toutes les trames à partir du *N_B(R)* reçu, et *N_A(S)* devient égal à *N_B(R)*.

Configurer les PVC

Dans notre réseau, un site peut communiquer avec plusieurs autres sites : les communications sont dites mutipoints. L'utilisation des interfaces non numérotées (*unnumbered*) est toujours possible, mais il vaut mieux considérer le réseau Frame Relay comme étant un réseau IP à part entière : il sera plus évolutif et pourra être administré. Notre plan d'adressage (voir chapitre 7) prévoit d'affecter le réseau 172.16.0.0/16.

De plus, l'interface du routeur doit être configurée avec plusieurs DLCI (un PVC par site distant). À Paris, notre opérateur nous a affecté les DLCI 40 et 60, respectivement pour Toulouse et Strasbourg :

Figure 10-9.
Exemples de DLCI



```
interface serial1
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay interface-dlci 40 broadcast
frame-relay interface-dlci 60 broadcast
```

Interface multipoint : jusqu'à 6 sites sur ce subnet

L'option "broadcast" indique que les broadcasts IP seront transmis sur la ligne série, permettant ainsi aux protocoles de routage tels qu'OSPF de fonctionner (cf. chapitre 11).

Correspondance entre adresses IP et DLCI

Sur le réseau intersite, les routeurs utilisent le protocole **Inverse ARP** pour découvrir les adresses IP des routeurs distants et les associer aux DLCI.

LE POINT SUR INVERSE ARP (RFC 1293)

Le protocole **ARP** (*Address Resolution Protocol* — RFC 826) permet à une station IP de connaître l'adresse physique (MAC ou autre) d'une autre station en connaissant son adresse IP. Le protocole **RARP** (*Reverse ARP* — RFC 903) permet à une station d'obtenir, à partir de son adresse MAC et auprès d'un serveur d'adresses, l'adresse IP qui lui a été affectée.

InARP (*Inverse ARP*) est le mécanisme inverse d'ARP : ce protocole permet à une station (typiquement un routeur) de connaître l'adresse IP du routeur se trouvant à l'autre bout d'un circuit virtuel (Frame Relay ou ATM). Les paquets envoyés sont identiques à ceux d'ARP. Le routeur envoie une requête InARP, et attend une réponse contenant l'adresse IP du routeur distant. Le routeur associe alors l'adresse IP reçue au DLCI local du circuit virtuel. Plusieurs routeurs distants peuvent répondre si les circuits virtuels sont de type multipoint. InARP fonctionne sur le même principe pour d'autres protocoles, tels que Decnet, Apple Talk ou IPX.

Si vous rencontrez des problèmes d'incompatibilité avec ce protocole, ou si InARP n'est pas supporté par le routeur distant, vous devez associer manuellement l'adresse IP de destination avec le DLCI :

```
interface serial 1
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay map ip 172.16.0.6 40 broadcast
frame-relay map ip 172.16.0.2 60 broadcast
```

L'opérateur configure les mêmes DLCI sur son équipement (FRAD ou commutateur) relié au routeur.

Configurer les SVC

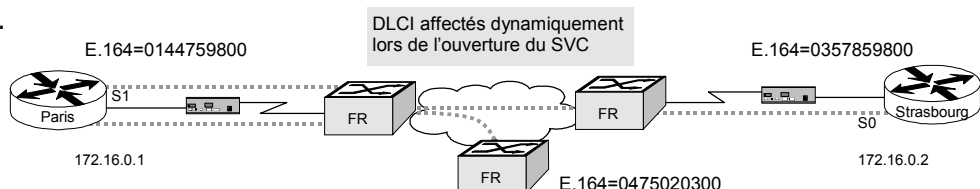
Le choix de circuits virtuels commutés peut être intéressant à double titre :

- Du point de vue du coût, tout d'abord. L'opérateur facture le SVC uniquement lorsqu'il est ouvert (facturation à la durée et/ou au volume).
- Du point de vue des performances ensuite. Si de nombreux CV doivent être utilisés (cas d'un réseau parfaitement maillé, par exemple), fermer ceux qui sont inutiles permet de libérer des ressources mémoire et CPU dans les routeurs et les commutateurs.

En plus de la signalisation Q.933 et de la procédure LAP-F, les SVC nécessitent l'emploi d'un troisième mécanisme, celui de **l'adressage** (voir la fin de ce chapitre pour plus de détails). Frame-Relay utilise soit celui du RNIS (E.164) dans le cas des réseaux publics, soit X.121 dans le cas des réseaux privés et publics. Cet adressage de niveau 3 est utilisé par le protocole de signalisation Q.933 pour ouvrir les SVC, c'est-à-dire affecter dynamiquement les DLCI d'entrée et de sortie dans chaque commutateur traversé. Les trames sont ensuite acheminées en fonction des DLCI qui réalisent un adressage de niveau 2.

Cette fois, l'activation des SVC nécessite une configuration manuelle, non des DLCI, mais des adresses E.164 ou X.121, ainsi que l'association des SVC avec les adresses IP.

Figure 10-10.
*Affectation
dynamique
des DLCI.*



```

interface ser 1
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay svc
frame-relay lmi-type q933a
map group strasbourg
map group toulouse

map-list strasbourg source-addr E.164 0144759800 dest-addr E.164
0357859800
ip 172.16.0.2 class operateur broadcast ietf

map-list toulouse source-addr E.164 0144759800 dest-addr E.164 0475020300
ip 10.16.0.6 class operateur broadcast ietf

map-class frame-relay operateur
frame-relay traffic-rate 256000 384000

```

Les CIR et AIR sont utilisés lors de la négociation qui a lieu à l'ouverture du SVC. Ces paramètres doivent être identiques à ceux configurés dans le commutateur.

Gérer la qualité de service

Quel que soit le type de CV, le routeur peut moduler les flux de données selon l'état du réseau (en cas de congestion). Sur les routeurs Cisco, la qualité de service est gérée en activant la fonction *traffic shaping* (voir chapitre 14). À cette occasion, le routeur met en place une file d'attente par DLCI, et adapte le flux en fonction des informations envoyées par le commutateur :

```

interface serial 1
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay traffic-shaping

```

L'activation de la fonction **ELMI** (*Extended Local Management Interface*) permet au routeur de connaître automatiquement les paramètres Frame Relay en recevant les messages LMI du commutateur (sur le DLCI 0). Le commutateur communique ainsi au routeur les valeurs des paramètres CIR, Be et Bc :

```
frame-relay qos-autosense
```

Les commutateurs peuvent informer les routeurs de l'état du réseau de deux manières : via les messages **CLLM** (*Consolidated Link Layer Management*) ou via le bit **BECN** (*Backward Explicit Congestion Notification*). Nous préférons le mode CLLM (appelé *foresight* chez Cisco) au mode BECN, moins fiable (voir encadré "Le point sur Frame Relay").

```

interface s 1
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay traffic-shaping
frame-relay qos-autosense
frame-relay class opérateur
!
map-class frame-relay opérateur
frame-relay adaptive-shaping foresight

```

Pour les SVC, ajoutez :
 Frame relays svc
 map group
 map-list ...
 ip ... class opérateur ...

...et supprimez cette commande.

Il se peut que les fonctions CLLM et ELMI ne soient pas disponibles sur le FRAD de l'opérateur ou que l'on y rencontre des incompatibilités. La solution est donc de configurer manuellement tous les paramètres. L'équivalent de la configuration précédente est, de ce fait, plus complexe :

```

interface s 1
bandwidth 512
ip address 172.16.0.1 255.255.255.248
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay traffic-shaping
frame-relay class opérateur

map-class frame-relay opérateur
frame-relay traffic-rate 256000 384000

```

La commande *traffic-rate* indique le CIR et l'AIR, que nous avons respectivement positionnés à 256 Kbit/s et 384 Kbit/s. Si la valeur de l'AIR est omise, la valeur par défaut est celle du débit de la ligne indiquée par la commande *bandwidth*.

Il est possible de définir plus finement les paramètres de qualité de service décrits dans la norme Q.922. Le routeur Cisco permet même de le faire dans les deux sens, bien que, généralement, les opérateurs proposent des valeurs identiques afin de simplifier la configuration des commutateurs et la grille tarifaire. Le profil "personnalisé" présenté ici se substitue alors au profil opérateur précédemment décrit :

```

map-class frame-relay personnalise
frame-relay cir in 1280000
frame-relay bc in 256000
frame-relay be in 256000
frame-relay cir out 2560000
frame-relay bc out 256000
frame-relay be out 128000
frame-relay idle-timer 30

```

$CIR = Bc/Tc$

$EIR = Be/Tc$

$AIR = CIR + EIR = (Bc+Be) / Tc$

À partir des trois valeurs CIR, Bc et Be, on peut déduire celle de Tc (2 secondes en entrée et 1 seconde en sortie).

Toujours dans l'optique de gérer la qualité de service, vous pouvez affecter une priorité plus faible à certains protocoles en positionnant le bit DE dans les trames qui les véhiculent :

```
int s 1
frame-relay de-group (1) 50 ← Liste de DLCI (rien = tous les DLCI)

frame-relay de-list (1) protocol ip characteristic tcp 20
frame-relay de-list 1 interface e 0 characteristic list 100 access-list
101
...
```

Par exemple, les protocoles tels que FTP pourront être marqués de cette manière afin de privilégier la voix sur IP en cas de congestion du réseau.

Les sous-interfaces

Dans certains cas, il peut être intéressant d'utiliser le principe des sous-interfaces proposé par Cisco, afin d'activer un secours RNIS individuellement, par PVC, et non pas sur la chute de l'interface série. Les sous-interfaces permettent également de configurer le *traffic shaping*, ainsi que d'autres paramètres Frame Relay PVC par PVC.

Sur notre routeur Cisco, l'interface physique est configurée en Frame Relay et est associée à des sous-interfaces logiques (une par DLCI). Nous avons choisi un DLCI par site distant, donc un mode point à point pour chaque sous-interface, ce qui permet de ne pas utiliser d'adresse IP sur la liaison WAN.

```
interface s 1
encapsulation frame-relay ietf
frame-relay lmi-type q933a
frame-relay traffic-shaping
frame-relay qos-autosense

interface serial1.1 point-to-point
ip address 172.16.0.1 255.255.255.252
frame-relay interface-dlci 40
    class opérateur

interface serial1.2 point-to-point
ip address 172.16.0.5 255.255.255.252
frame-relay interface-dlci 60
    class personnalisé
```

Pour terminer, il est possible d'optimiser l'utilisation des liaisons série en compressant les données (norme FRF.9 du Frame Relay Forum).

`frame-relay map payload-compress frf9 stac` ← Méthode Stacker

Cette commande peut s'appliquer à une interface physique ou à une sous-interface.

Mettre en place un réseau ATM

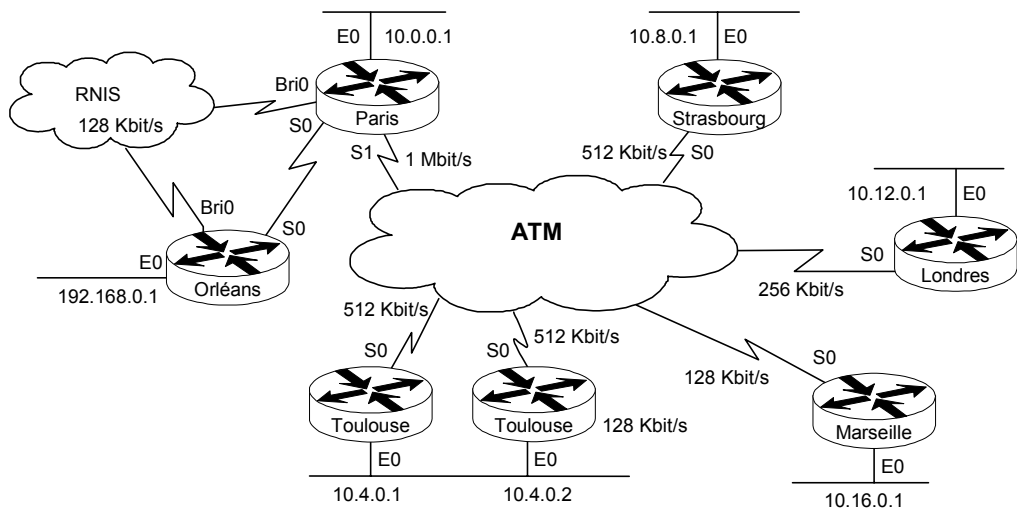
Les réseaux ATM sont réservés aux hauts débits : à partir de 34 Mbit/s en France, et de 45 Mbit/s aux États-Unis. La facturation associée est donc très élevée.

Nous nous plaçons ici dans cette situation.

Comme précédemment, l'opérateur local installe son modem numérique (le CSU) dans vos locaux. L'opérateur retenu pour le réseau ATM installe éventuellement un commutateur ATM qu'il connecte au modem. L'interface série du routeur sera ensuite connectée à l'un des ports du commutateur.

Figure 10-11.

Réseau intersite reposant sur un réseau opérateur ATM.



Qualité de service et facturation

Avec ATM, les opérateurs peuvent contrôler précisément la qualité de service. Ils élaborent leur offre commerciale à partir des mécanismes offerts par le protocole.

Par exemple, ATM permet de **gérer la congestion** du réseau et d'offrir différentes **classes de service**, ce qui permet de proposer une facturation adaptée à chaque type de situation (flux voix et/ou données et/ou vidéo). Comment ? La réponse est intimement liée au fonctionnement du protocole ATM.

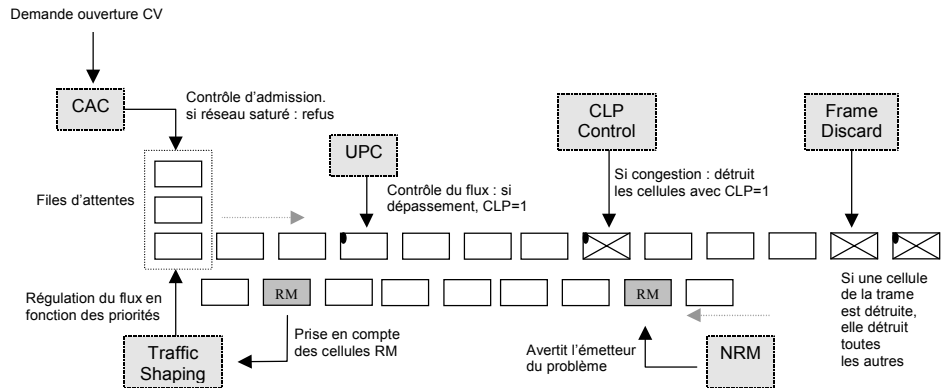
La gestion du trafic : TMS 4.0 (ATM Forum af-tm-0056.000)

La spécification TMS (*Traffic Management Specification*) définit les paramètres et procédures relatifs à la gestion du trafic ainsi qu'à la qualité de service. Elle reprend les normes de l'ITU en précisant leur fonctionnement (paramètres et procédures pris en charge ou non) et, surtout, en les modifiant. Elle définit également les procédures et algorithmes permettant aux commutateurs ATM d'assurer la qualité de service demandée.

Procédure	Description
CAC (<i>Connection Admission Control</i>)	Avant d'autoriser l'ouverture d'un CV, chaque commutateur vérifie qu'il pourra bien assurer la qualité de service demandée sans remettre en cause celles déjà accordées.
UPC (<i>Usage Parameter Control</i>)	Chaque commutateur vérifie que le flux émis dans un CV respecte bien la QOS demandée. En cas de dépassement, les cellules sont marquées avec le bit CLP positionné à "1" et peuvent être détruites.
CLP Control (<i>Cell Loss Priority control</i>)	L'équipement terminal (par exemple, le routeur) peut positionner le bit CLP à "1" dans les cellules qui ne sont pas prioritaires. En cas de congestion ou de dépassement de la QOS, les commutateurs détruiront ces cellules en priorité.
NRM (<i>Network Resource Management</i>)	Si des signes de congestion apparaissent (remplissage d'une file d'attente, par exemple), les commutateurs peuvent émettre des cellules RM à destination de l'émetteur d'un flux (retour sur la qualité de service - <i>feedback</i>) l'invitant à réguler son trafic via le mécanisme de <i>traffic shaping</i> .
Traffic Shaping	Les équipements terminaux et les commutateurs peuvent modifier les caractéristiques du flux émis (réduction du débit, régulation d'un trafic erratique, etc.).
Frame Discard	Si une cellule est détruite, le commutateur peut détruire toutes les autres cellules appartenant aux mêmes données (par exemple, toutes les cellules d'un paquet IP).

Figure 10-12.

Gestion
de la qualité
de service
par ATM.



Les classes de service ATM Transfer Capabilities (ITU I.371 et TMS)

Les procédures détaillées ci-dessus ont pour but d'assurer une certaine qualité de service aux applications. Celles-ci ont le choix entre **cinq classes de service** adaptées à différents types de flux. Par exemple, les cellules véhiculant un canal voix d'une conversation téléphonique doivent être transmises avec une grande régularité (les horloges de l'émetteur et du récepteur doivent être synchronisées). En revanche, un flux de type réseau local, qui est par nature erratique (trafic par rafales), ne nécessite pas la même qualité de service.

Ces classes de service sont accessibles via différentes interfaces d'accès à ATM, appelées **AAL** (*ATM Adaptation Layer*).

Classe de service	Caractéristiques du trafic	Applications
CBR (<i>Constant Bit Rate</i>)	Le débit est constant et garanti.	Voix non compressée en émulation de circuit (accès de préférence via AAL-1)
rt-VBR (<i>Real-Time Variable bit Rate</i>)	Le débit est variable et est garanti. Le délai de transit est garanti et varie peu.	Données (accès de préférence via AAL-2)
nrt-VBR (<i>Non Real-Time Variable bit Rate</i>)	Le débit est variable et est garanti.	Vidéo + voix compressées (accès de préférence via AAL-2 ou AAL-5)
ABR (<i>Available Bit Rate</i>)	Le débit est variable et peut être modifié si le réseau le demande (en cas de congestion, par exemple).	Données (accès de préférence via AAL-5)
UBR (<i>Unavaible Bit rate</i>)	Acheminement sans garanti, au mieux des capacités du réseau (<i>best effort</i>).	Données (accès de préférence via AAL-5)

La classe de service UBR n'offre aucune garantie de service, tandis que l'ABR est la plus utilisée, notamment par l'UNI 4.0.

LA SIGNALISATION UNI 4.0 (ATM FORUM AF-SIG-0061.000)

L'UNI (*User Network Interface*) reprend les normes de l'ITU en précisant leur fonctionnement (paramètres et procédures pris en charge ou non), et surtout en les modifiant :

- Appels point à point (**Q.2931**) sans les fonctions OAM (*Operations, Administration and Maintenance*).
- Adressage **NSAP** (ISO 8348, ITU X.213, RFC 1629).
- Couche **SAAL** (Q.2100, Q.2110, Q.2130) ; vpi/vci = 0/5.
- Appels point à multipoint (**Q.2971**) avec extensions pour que les feuilles puissent accepter de nouveaux participants à la conférence en cours (dans Q.2971, une conférence à plusieurs est gérée sous forme d'arbre, dont seule la racine permet à d'autres utilisateurs de rejoindre une conférence).
- Direct Dialling In (Q.2951).

L'ATM Forum a ajouté les extensions suivantes à Q.2931 :

- Gestion des **adresses de groupe** ATM (adresses *anycast*) identifiées par le préfixe C50079. Par exemple, le groupe C50079.00000000000000000000.00A03E000001.00 permet de joindre le LECS (*LAN Emulation Configuration Server*) .
- Négociation des caractéristiques de la connexion (CBR, VBR, rt-VBR, ABR, UBR).
- Présentation individuelle des paramètres de qualité de service.
- Paramètres décrivant la classe de service **ABR**.
- **Signalisation proxy** : un nœud ATM peut gérer la signalisation à la place d'un autre nœud qui ne la supporte pas. Permet également à un serveur disposant de plusieurs interfaces ATM de ne posséder qu'une seule adresse NSAP.

En l'état actuel de la norme, les cellules RM (**Resource Management**) ne sont utilisées que pour réguler le trafic ABR.

Connecter le routeur au réseau de transport

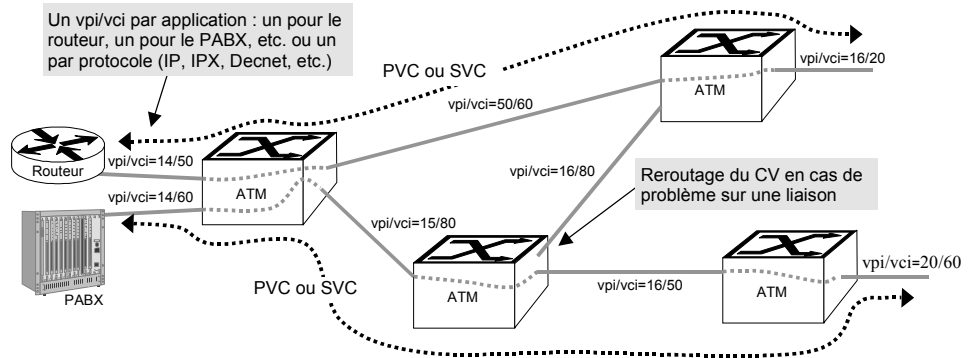
La norme ATM définit un ensemble de protocoles reposant sur la commutation de petites cellules de taille fixe (58 octets dont 5 d'en-tête).

Chaque cellule est identifiée par un couple d'adresses **VPI** (*Virtual Path Identifier*) et **VCI** (*Virtual Channel Identifier*), adresse locale partagée par deux commutateurs. Il n'y a pas d'adressage de bout en bout, mais uniquement un adressage point à un point entre deux commutateurs. Un commutateur recevant une trame avec un VPI/VCI donné la routera sur un autre port et l'enverra avec un autre VPI/VCI, et ainsi de suite.

La connexion entre deux commutateurs ATM s'effectue par l'ouverture de circuits virtuels permanents (**PVC**, *Permanent Virtual Circuit*) ou commutés (**SVC**, *Switched Virtual Circuit*), c'est-à-dire ouverts à la demande via le protocole de signalisation Q.2931 (qui utilise quelques Kbit/s du VPI/VCI = 0/5).

L'intérêt des SVC est que la qualité de service peut être spécifiée à la demande permettant ainsi de réduire (encore) les coûts. Sur un PVC, le débit est fixé une fois pour toutes et engendre un coût fixe.

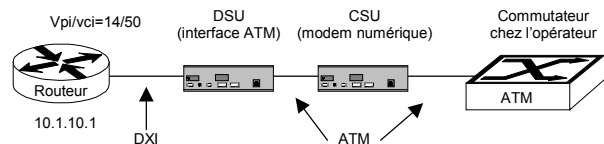
Figure 10-13.
Circuits virtuels et vpi/vci ATM.



Si le routeur ne dispose pas d'interface ATM

La manière la plus simple de connecter nos réseaux locaux est de configurer l'interface série du routeur en mode **DXI** (*Data eXchange Interface*). Ce mode de fonctionnement implique l'utilisation d'un DSU (*Data Service Unit*) externe fourni par l'opérateur. Le DSU se charge d'adapter le flux série émis et reçu par le routeur en cellules ATM.

Figure 10-14.
Connexion d'un routeur à un commutateur ATM.



L'exemple suivant montre un PVC configuré sur notre routeur parisien à destination du réseau de Strasbourg :

```
interface serial 0
ip address 172.16.0.1 255.255.255.248
encapsulation atm-dxi
dxi pvc 14 50 mux
dxi map ip 172.16.0.2 14 50 broadcast
```

VPI = 14 / VCI = 50

Le paramètre “mux” indique que nous avons choisi l’encapsulation de type multiplexage par circuit virtuel : un seul protocole (IP dans notre cas) utilisera le PVC identifié par le vpi 14 et le vci 50. L’encapsulation LLC/SNAP (paramètre “snap” à la place de “mux”) n’apporterait aucun avantage et ajouterait un overhead de 8 octets par paquet IP.

L’INTERFACE DXI (ATM FORUM AF-DXI-0014.000)

La norme **DXI** (*Data eXchange Interface*) définit le protocole de niveau 2 utilisé pour échanger les données entre un équipement non ATM (*via une interface série V.35 ou HSSI*) et un DSU ATM. Trois modes de fonctionnement sont possibles.

Dans le **mode 1a**, le routeur place ses données dans une trame DXI (SDU de 9 232 octets), puis le DSU réalise l’encapsulation AAL-5, la segmentation SAR AAL-5 et l’accès l’UNI. Ce mode supporte 1 023 CV simultanés.

Dans le **mode 1b**, le routeur réalise l’encapsulation AAL-3/4 puis place le résultat dans une trame DXI (SDU de 9 224 octets). Le DSU réalise la segmentation SAR AAL-3/4 et l’accès UNI. Le mode 1a est également supporté, mais le mode 1b doit être utilisé pour au moins 1 CV. Ce mode supporte 1 023 CV simultanés.

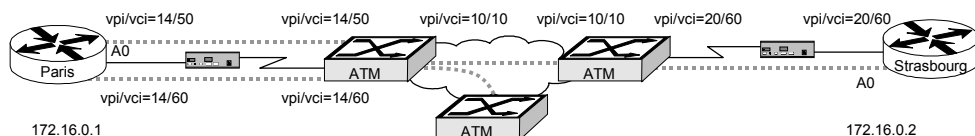
Dans le **mode 2**, le routeur réalise l’encapsulation AAL-3/4, puis place le résultat dans une trame DXI (SDU de 65 535 octets). En cas de besoin, le DSU peut convertir l’encapsulation AAL-3/4 en AAL-5 (changement d’encapsulation). Le DSU réalise ensuite l’assemblage SAR (AAL-3/4 ou AAL-5) et l’accès UNI. Ce mode supporte 16 millions de CV simultanés.

Si le routeur supporte ATM

La seconde solution consiste à insérer une carte ATM dans le routeur. Fonctionnellement, celle-ci se comporte comme un DSU et peut être directement connectée au CSU (c’est-à-dire le modem numérique) de l’opérateur ou à un commutateur.

Pour la même interconnexion Paris-Strasbourg et Paris-Toulouse, la configuration devient la suivante :

Figure 10-15.
Configuration
des vpi/vci
ATM



```

interface atm 0
ip address 172.16.0.1 255.255.255.248
atm pvc 1 14 50 aal5mux ip
atm pvc 2 14 60 aal5mux ip
map-group operateur
!
map-list operateur
ip 172.16.0.2 atm-vc 1 broadcast
ip 172.16.0.6 atm-vc 2 broadcast

```

Identifiant du PVC propre à Cisco
(interne au routeur)

L'encapsulation choisie, "aal5mux", est la même que précédemment : un CV transporte le protocole IP. Cependant, l'encapsulation LLC/SNAP, utilisée conjointement avec l'interface ATM, nous donne la possibilité supplémentaire de pouvoir exécuter **Inverse ARP**, et donc de simplifier la configuration :

```

interface atm 0
ip address 172.16.0.1 255.255.255.248
atm pvc 1 14 50 aal5snap inarp 5
atm pvc 2 14 60 aal5snap inarp 5

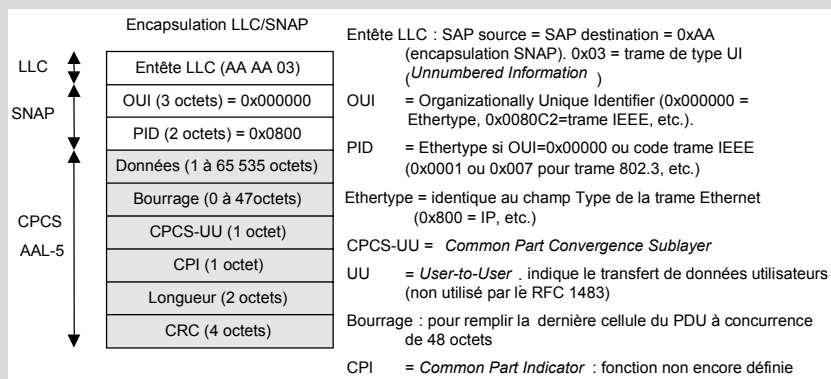
```

La valeur 5 affectée à "inarp" indique que la procédure Inverse ARP est lancée toutes les 5 minutes. La correspondance manuelle entre adresse IP et circuit virtuel n'est donc plus nécessaire.

L'ENCAPSULATION DES PROTOCOLES DANS LES CELLULES ATM (RFC 1483)

Le transport d'un protocole dans des cellules ATM requiert l'utilisation de l'interface **AAL-5** (*ATM Adaptation Layer 5*), accessible via la couche **CPCS** (*Common Part Convergence Sublayer*). L'encapsulation peut être réalisée de deux manières : soit en dehors d'ATM, soit au niveau d'ATM.

La première méthode permet de transporter plusieurs protocoles dans un seul circuit virtuel ATM ; elle utilise pour cela l'encapsulation **LLC/SNAP** (*Logical Link Control/Sub Network Access Protocol*).



La seconde méthode, appelée **multiplexage par circuit virtuel**, consiste à encapsuler le protocole directement dans le PDU AAL-5, ce qui implique l'utilisation d'un circuit virtuel ATM par protocole.

Configurer les SVC

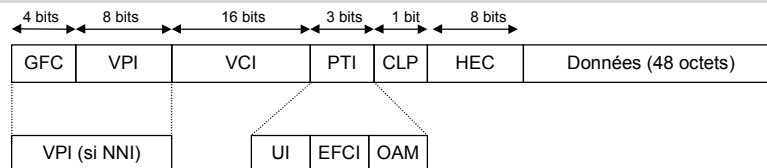
Si l'opérateur nous en donne la possibilité, la mise en place de SVC permet, comme pour Frame Relay, de réaliser des économies : l'opérateur facture en fonction de l'utilisation (à la durée et/ou au volume). De plus, si les CV sont nombreux, la fermeture des SVC lorsqu'ils ne sont pas utilisés permet de libérer des ressources CPU et mémoire dans les routeurs et les commutateurs.

L'utilisation des SVC implique l'activation de deux nouveaux mécanismes :

- Un PVC pour le protocole SAAL (*Signaling ATM Adaptation Layer*) qui gère les SVC (ouverture, fermeture, etc.). Le vpi/vci utilisé est 0/5.
- L'utilisation d'un adressage global (de niveau 3) permettant d'identifier les nœuds du réseau. L'adressage utilisé par ATM est de type NSAP ; trois encapsulations d'adresses sont possibles : DCC, ICD et E.164 (reportez-vous à la fin de ce chapitre pour plus de détails).

LE POINT SUR ATM (ITU I.361)

ATM (*Asynchronous Transfer Mode*) découpe la bande passante en tranches de temps fixe appelées **cellules**.



GFC (*Generic Flow Control*). Priorité de la cellule (0 = la plus basse).

VPI (*Virtual Path Identifier*). Identifie le chemin virtuel (255 possibilités).

VCI (*Virtual Channel Identifier*). Identifie la voie virtuelle au sein du chemin virtuel (65 535 possibilités).

PTI (*Payload Type Indicator*). Le premier bit indique si la cellule transporte des données de contrôle ou des données utilisateur. Dans ce dernier cas, le deuxième bit, appelé **EFCI** (*Explicit Forward Congestion Indication*), indique à l'application qu'il faut prévoir des délais d'acheminement pour les cellules à venir (suite à une congestion par exemple). Positionné à "1", le troisième bit indique que le champ d'information contient des données utilisées par les applications d'administration **OAM** (*Operations, Administration and Maintenance*).

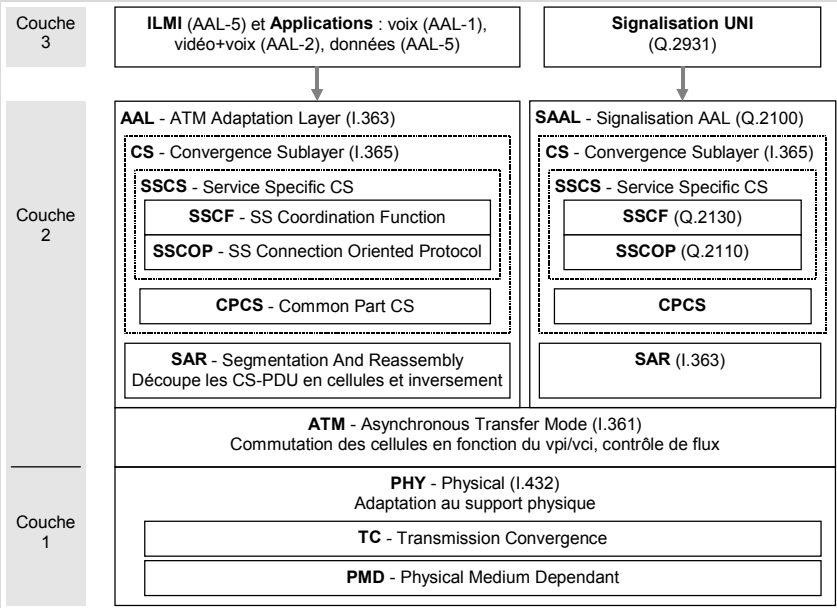
CLP (*Cell loss Priority*). Positionné à "1", ce bit indique que la cellule peut être détruite par le commutateur en cas de congestion.

HEC (*Header Error Control*). Cet octet permet à la couche TC (*Transmission Convergence*) d'opérer un contrôle d'erreur sur l'en-tête de la cellule.

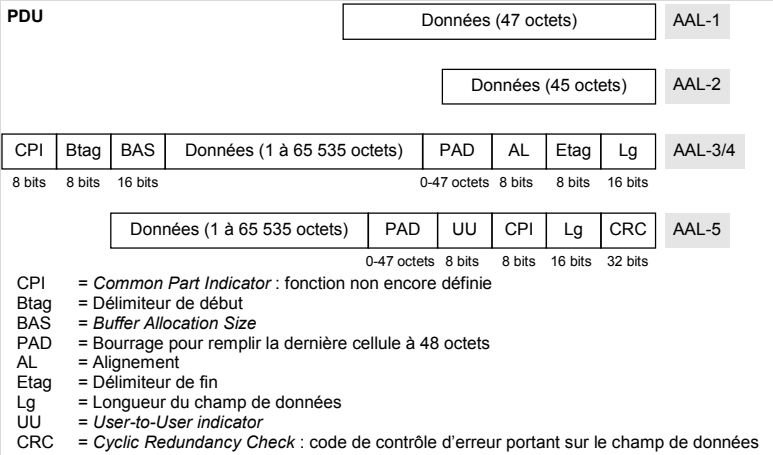
...

LE POINT SUR ATM (ITU I.361 – SUITE)

Les applications transmettent leurs données à la couche **AAL** (*ATM Adaptation Layer*) qui se charge de les convertir en cellules, puis de les envoyer en respectant le niveau de service demandé (AAL-1 à AAL-5 et SAAL).

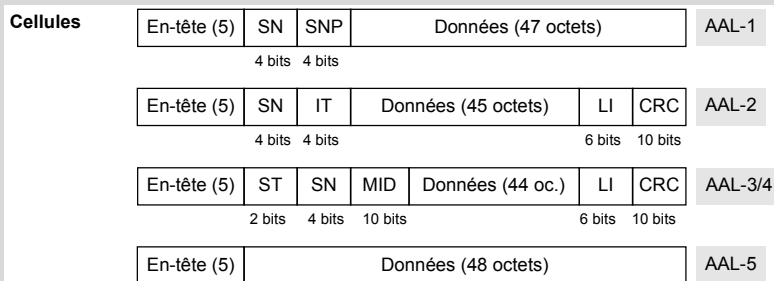


La couche AAL, et donc les sous-couches qui la composent, est adaptée à chaque type de service (AAL-1, AAL-2, AAL-3/4, AAL-5 et SAAL). Par exemple, la couche **CS** accepte des données au format **CS-PDU** (*Convergence Sublayer-Protocol Data Unit*).



LE POINT SUR ATM (FIN)

De même, la couche **SAR** structure différemment les 48 octets du champ de données des cellules ATM (appelé SAR-PDU).



SN = *Sequence Number* : détecte les cellules manquantes ou erronées
 SNP = *Sequence Number Protection* : code autocorrecteur portant sur le SN
 IT = *Information Type* : début, continuation ou fin d'un CS-PDU
 ST = *Segment Type* : début, fin, continuation ou segment simple
 MID = *Multiplexing Identifier* : partage d'un circuit virtuel par plusieurs applications de la couche SAR
 LI = *Length Indicator* : nombre d'octets significatifs dans le cas d'une cellule partiellement remplie
 CRC = *Cyclic Redundancy Check* : code de contrôle d'erreur portant sur le champ de données

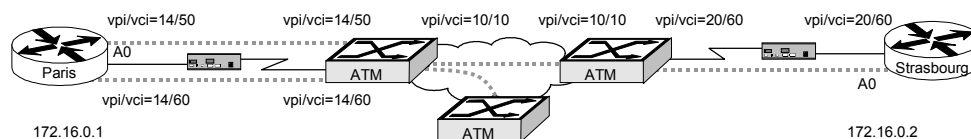
En comparant les deux schémas précédents, on constate que les PDU des AAL-1 et AAL-2 sont directement insérés dans une cellule ATM. La couche AAL n'utilise donc pas obligatoirement les mécanismes de segmentation et d'assemblage.

Il existe également une couche SSCS pour Frame Relay (I.365.1).

La commande **map-list** permet de configurer manuellement la correspondance entre adresses NSAP et adresses IP :

Figure 10-16.

Configuration
des vpi/vci
ATM



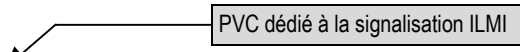
```
interface atm 0
ip address 172.16.0.2 255.255.255.248
atm nsap-address
  47.0091.81.000000.0061.705b.7701.0800.200c.1A2B.01
AFI|ICD|Préfixe DSP|ESI|SEL
atm pvc 1 0 5 qsaal
map-group operateur
!
map-list operateur
ip 172.16.0.2 atm-nsap 47.0091.81.000000.0061.705b.7701.0800.200c.1000.02
broadcast
```

PVC dédié à la signalisation
Q.2931

Adresses ATM

```
ip 172.16.0.6 atm-nsap 47.0091.81.000000.0061.705b.88a7.0900.20ab.2000.01
broadcast
```

Une autre manière d'affecter l'adresse NSAP à l'interface est d'utiliser la signalisation **ILMI** (*Integrated Local Management Interface*). Elle permet au routeur d'obtenir le préfixe de l'adresse, le DSP (champs ESI et SEL) étant toujours affecté par l'équipement terminal (le routeur). Le PVC dédié à ILMI est : vpi/vci = 0/16 :



```
atm pvc 2 0 16 ilmi
atm esi-address 0800200c1000.02
```

ESI		SEL
-----	--	-----

La première partie du chiffre correspond au champ *ESI* (6 octets, l'équivalent d'une adresse MAC) ; la seconde au champ *Selector* (1 octet) de l'adresse NSAP. Le reste de l'adresse (13 octets) est affecté automatiquement par le commutateur ATM *via* ILMI.

L'activation de cette signalisation permet également au routeur et au commutateur de surveiller l'état des circuits virtuels.

LA SIGNALISATION ILMI 4.0 (ATM FORUM AF-ILMI-0065.00)

La signalisation ILMI (*Integrated Local Management Interface*)* est un protocole permettant aux commutateurs ATM d'échanger des informations avec les équipements terminaux (stations ATM, routeurs, etc.), telles que :

- la configuration et l'état des circuits virtuels ;
- les préfixes des adresses NSAP ;
- les services et options supportés par les équipements.

Le protocole utilisé par ILMI est **SNMP** (*Simple Network Management Protocol* — RFC 1157), qui utilise les services de **AAL-5** (vpi/vci = 0/16). La base de données **MIB** (*Management Information Base*) contient des informations relatives aux couches physique et ATM, aux circuits virtuels (état des vpi/vci), aux adresses, ainsi qu'aux services supportés et pouvant être négociés (version UNI, nombre maximal de vpi/vci, qualité de service, etc.).

Par exemple, ILMI permet de configurer les LEC (*LAN Emulation Client*) et de trouver le LECS (*LAN Emulation Configuration Server*).

La bande passante utilisée par ILMI ne doit pas dépasser 1 % du débit de la ligne, et 5 % en pic.

*Note : le 1^{er} « I » signifiait *Interim* car l'ATM Forum attendait la normalisation de l'ITU. Mais celle-ci s'étant fait attendre, l'ATM Forum a définitivement entériné sa proposition de norme.

Gérer la qualité de service

Les normes prévoient de nombreux et complexes mécanismes pour gérer la qualité de service sur ATM. Seuls certains d'entre eux sont disponibles sur nos routeurs.

Pour les PVC, seules les caractéristiques du flot de données peuvent être spécifiées :

```
atm pvc 1 14 50 aal5snap 384 256 inarp 5
```

Cette commande précise qu'un maximum de 384 Kbit/s sera alloué à notre PVC et que le débit moyen du trafic sera de 256 Kbit/s.

Les possibilités de paramétrage sont plus étendues en ce qui concerne les SVC puisque l'on peut demander l'activation de **classes de service**. Celles-ci sont décrites implicitement par des combinaisons de paramètres :

```
map-list operateur
ip 172.16.0.2 atm-nsap 47.0091.81.000000.0061.705b.7701.0800.200c.1000.02
broadcast class traficUBR
ip 172.16.0.6 atm-nsap 47.0091.81.000000.0061.705b.88a7.0900.20ab.2000.01
broadcast class traficNrtVBR
```

```
map-class traficUBR
atm forward-peak-cell-rate-clp1 384
atm backward-peak-cell-rate-clp1 256
```

Cette combinaison de paramètres active la classe de service nrt-VBR.

```
map-class traficNrtVBR
atm forward-peak-cell-rate-clp1 384
atm forward-sustainable-cell-rate-clp1 256
atm forward-max-burst-size 128
atm backward-peak-cell-rate-clp1 384
atm backward-sustainable-cell-rate-clp1 256
atm backward-max-burst-size 128
```

La qualité de service peut être différente dans chaque sens.

Le suffixe "clp1" indique que le paramètre s'applique aux cellules dont le bit CLP est positionné à 1 ou 0. Le suffixe "clp0" permet d'appliquer les mêmes paramètres aux cellules dont le bit CLP est à 0 (c'est-à-dire non marquées en suppression).

Dans le cas de la classe de service "traficUBR", aucune bande passante n'est réservée pour le SVC. C'est le mode de fonctionnement par défaut si aucun paramètre n'est spécifié.

Enfin, il est possible d'activer la procédure de contrôle CAC (*Connection Admission Control*) au niveau du routeur.

```
atm sig-traffic-shaping strict
```

Avec cette commande, l'ouverture d'un SVC ne sera possible que si les commutateurs ATM sont capables d'assurer la qualité de service demandée.

Les paramètres décrivant les classes de service (ITU I.356 et ATM Forum TMS 4.0)

Chaque classe de service (UBR, ABR, etc.) est définie par un ensemble de paramètres qui décrivent les caractéristiques du flux qui sera généré, ainsi que la qualité de service demandée. Ces paramètres peuvent, par exemple, être configurés sur nos routeurs.

Paramètre	Description du trafic
PCR (<i>Peak Cell Rate</i>)	Débit maximal autorisé en pointe (nombre maximal de cellules par secondes)
SCR (<i>Sustainable Cell Rate</i>)	Débit moyen autorisé (nombre moyen de cellules par secondes)
MBS (<i>Maximum Burst Size</i>)	Nombre de cellules autorisées pendant le débit en pointe. PCR/MBS = durée pendant laquelle le débit en point est autorisé
MCR (<i>Minimum Cell Rate</i>)	Débit minimal demandé (nombre minimal de cellules par secondes)
Paramètre	Description de la qualité de service
CDV (<i>Cell Delay Variation</i>)	Demande d'une variation maximale du délai de transit (la gigue - <i>jitter</i>). Le flux doit être le plus constant possible ($\pm$ CDV millisecondes). La fonction UPC utilise pour cela l'algorithme GCRA (<i>Generic Cell Rate Algorithm</i>) de type <i>Leaky Bucket</i> ⁽¹⁾ .
MCTD (<i>Maximum Cell Transfer Delay</i>)	Délai maximal de transit des cellules entre l'UNI de l'émetteur et l'UNI du récepteur.
CLR (<i>Cell Loss Ratio</i>)	Pourcentage acceptable de cellules pouvant être perdues (appliqué aux cellules ayant le bit CLP positionné à "0")
Autre paramètre	Description
RM (<i>Resource Management</i>)	Traitement des cellules RM permettant d'adapter le trafic en fonction de l'état du réseau reporté par les commutateurs (<i>feedback</i>).

(1) *Leaky Bucket* signifie littéralement "seau (d'eau) percé". Voir le chapitre 14 à ce sujet.

Le tableau suivant indique les combinaisons autorisées. Ainsi, le paramètre PCR seul indique implicitement la classe de service UBR. En revanche, la configuration des paramètres SCR et MBS seuls ne correspond à aucune classe de service ; elle représente donc une combinaison invalide.

Classe de service	Description du trafic				Qualité de service			Autre paramètre
	PCR	SCR	MBS	MCR	CDV	MCTD	CLR	
CBR	X				X	X	X	
rt-VBR	X	X	X		X	X	X	
nrt-VBR	X	X	X		X		X	
ABR	X			X	X		X	X
UBR	X							

D'autres paramètres relatifs à la gestion de la qualité de service sont prévus par la norme :

- CER (*Cell Error Ratio*). Taux maximal de cellules pouvant être en erreur.
- SECBR (*Severely Errored Cell Block Ratio*). Taux maximal de cellules consécutives (par blocs de N définis dans ITU I.610) pouvant être en erreur.
- CMR (*Cell Misinsertion Rate*). Taux maximal de cellules pouvant être mal insérées (c'est-à-dire dont les erreurs portant sur l'en-tête n'ont pas été détectées).

La détection des erreurs et la surveillance des performances sont réalisées par la couche OAM (*Operations, Administration and Maintenance* – ITU I.610).

L'adressage

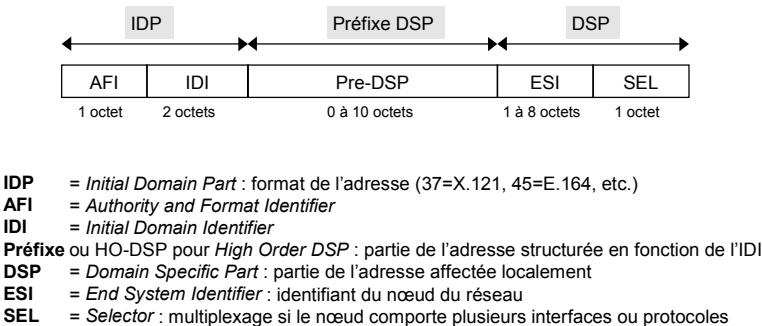
L'adressage NSAP (ISO 8348, ITU X.213, RFC 1629)

NSAP (*Network Service Access Point*) définit un **adressage global de niveau 3** permettant d'identifier les utilisateurs d'un réseau ATM, Frame Relay ou RNIS. Dans le cas du RNIS, il s'agit tout simplement du numéro de téléphone.

	Adresse locale	Adresse globale
Protocole	Niveau 2	Niveau 3
Frame Relay	DLCI	E.164 ou X.121
ATM	VPI/VCI	NSAP (encapsulation DCC, ICD ou E.164) ou E.164
RNIS	SAPI/TEI	E.164

L'apparent paradoxe d'un adressage de niveau 3 utilisé par des protocoles de niveau 2 s'explique par le fait que ces protocoles agissent en **mode connecté**. L'adressage global est utilisé par les **protocoles de signalisation** (situés au niveau 3) pour établir les communications de niveau 2 (les circuits commutés RNIS ou virtuels ATM) ; la commutation de circuit RNIS ou de cellule ATM n'utilise ensuite que des adresses locales de niveau 2.

Figure 10-17.
Format générique
d'une adresse NSAP.

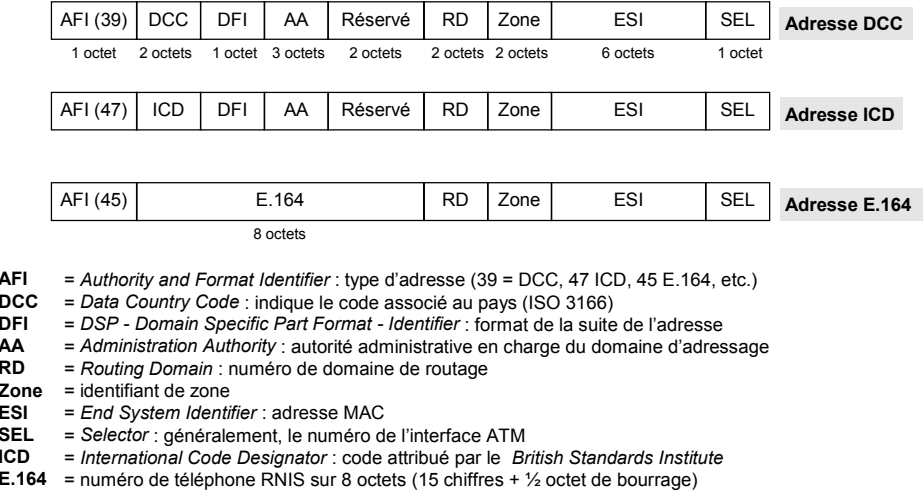


L'adressage ATM

ATM peut utiliser trois types d'adresses encapsulées au format NSAP :

- DCC pour les réseaux publics et privés ;
- ICD pour les réseaux privés ;
- E.164 pour les réseaux publics, soit de manière native, soit encapsulée dans une adresse NSAP.

Figure 10-18.
Format
des adresses ATM.



Dans les réseaux publics, ATM utilise une adresse E.164 native.

L’adressage Frame Relay

Les réseaux Frame Relay peuvent utiliser deux types d’adresses :

- **E.164**, identique aux numéros de téléphones RNIS
- **X.121**, identique aux réseaux X.25.

Figure 10-19.
Format
d’une adresse E.164.

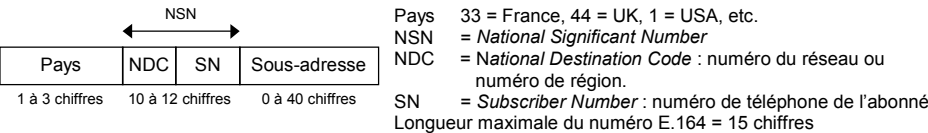
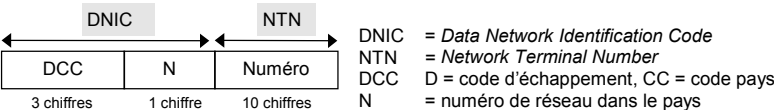


Figure 10-20.
Format
d’une adresse X.121.



Interopérabilité entre Frame Relay et ATM

Les réseaux Frame Relay et ATM peuvent être utilisés conjointement de deux manières : soit par encapsulation de Frame Relay dans les cellules ATM, soit par conversion de protocole. L’opérateur peut ainsi offrir un réseau avec une interface Frame Relay tout en utilisant ATM au sein de son réseau fédérateur.

La seconde méthode consiste à convertir une trame Frame Relay en un PDU AAL-5.

Frame Relay	ATM
Encapsulation NLPID ou SNAP (RFC 2427)	Encapsulation SNAP/LLC sur AAL-5 (RFC 1483)
Bit DE	Bit CLP
Bit FECN	Bit EFCI
LMI	ILMI
Inverse ARP (PVC)	Classical IP : Inverse ARP (PVC) et serveur ATM (SVC)

Assembler les briques du LAN et du WAN

Jusqu'à présent, nous avons utilisé différentes technologies, les unes adaptées aux réseaux locaux, les autres aux réseaux étendus.

Arrive un moment où les deux mondes doivent se rencontrer puisque la vocation des réseaux est de relier des hommes, qu'ils fassent ou non partie de la même entreprise.

Les réseaux locaux ont de plus en plus tendance à s'étendre au-delà d'un simple site pour former un réseau de campus, repoussant ainsi la frontière qui les sépare des réseaux étendus.

Dans ce chapitre vous apprendrez :

- à étendre le réseau fédérateur jusqu'au campus ;
- à configurer des VLAN ;
- à établir le lien entre commutateurs LAN et routeurs WAN ;
- à élaborer un plan de routage ;
- à configurer les protocoles de routage.

Mettre en place un réseau fédérateur

Les données du problème

Le réseau de 800 postes dont nous avons décrit l'installation au chapitre 6 fonctionne parfaitement. Or, voici que l'ouverture d'un nouvel immeuble à proximité est annoncée. Elle implique un changement d'échelle, puisqu'il s'agit d'une tour de quinze étages, représentant environ 2 000 connexions, à raison de 130 par étage en tenant compte des postes de travail, des imprimantes, des serveurs, etc.

Le câblage a été conçu en fonction des besoins potentiels en matière d'architecture, incluant à la fois la téléphonie (la voix), le réseau local (les données), et la diffusion vidéo (l'image). Les principes sont ceux qui ont été étudiés au chapitre 5.

La démarche

Il semble tout d'abord évident qu'il faudra au moins un réseau local par étage, afin de contrôler les flux, et sans doute plus, car il faut toujours s'attendre à des besoins spécifiques pour une population de 1 500 utilisateurs. Il est donc sage de prévoir une quarantaine de réseaux.

Un constat s'impose : s'il faut descendre près de quinze réseaux en *collapse backbone*, les équipements fédérateurs doivent disposer d'une très grande capacité. De plus, un réseau redondant est absolument nécessaire pour assurer une bonne qualité de service. En effet, à une telle échelle, un problème survient nécessairement quelque part (en vertu d'un principe de probabilité).

Le point central de l'architecture concerne donc les caractéristiques du réseau fédérateur pour lequel nous nous posons les questions suivantes :

- Quelle technologie ?
- Quels équipements ?
- Routeurs ou commutateurs de niveau 3 ?

Quelle technologie ?

Nous avons ici le choix entre Ethernet et ATM, sujet que nous avons abordé au cours du chapitre précédent.

Bien qu'adapté aux réseaux WAN, les constructeurs nous proposent d'utiliser ATM également pour les réseaux locaux. Choix étrange, car l'utilisation de ce protocole pose un certain nombre de problèmes :

- Il faut mettre en place une mécanique complexe pour adapter un réseau multipoint tel qu'Ethernet à un réseau ne fonctionnant qu'avec des circuits virtuels point à point.
- Il faut mettre en place une mécanique non moins complexe pour adapter les VLAN Ethernet au monde ATM.
- Le débit d'ATM est aujourd'hui limité à 622 Mbit/s, 155 Mbit/s étant le débit le plus fréquemment rencontré dans les entreprises. Face au Gigabit Ethernet, l'argument est donc mince.

Un certain nombre de standards permettent d'effectuer cette intégration. Il s'agit de LANE (*LAN Emulation*) pour les VLAN, de MPOA (*Multi Protocol Over ATM*) pour le routage et l'interconnexion des ELAN (*Emulated LAN*) et de Classical IP pour la correspondance entre adresses ATM et adresses IP.

Face à cela, Ethernet nous offre la simplicité et une panoplie de solutions homogènes et évolutives, du 10 mégabits au Gigabit. Nous choisirons donc cette technologie pour l'ensemble de notre réseau local, du poste de travail au réseau fédérateur.

Quels équipements ?

Le réseau fédérateur concentre tous les flux entre les réseaux d'étage d'une part, et entre ces derniers et les ressources communes d'autre part. Cela suppose que la majorité des flux est émise entre les utilisateurs d'un étage donné.

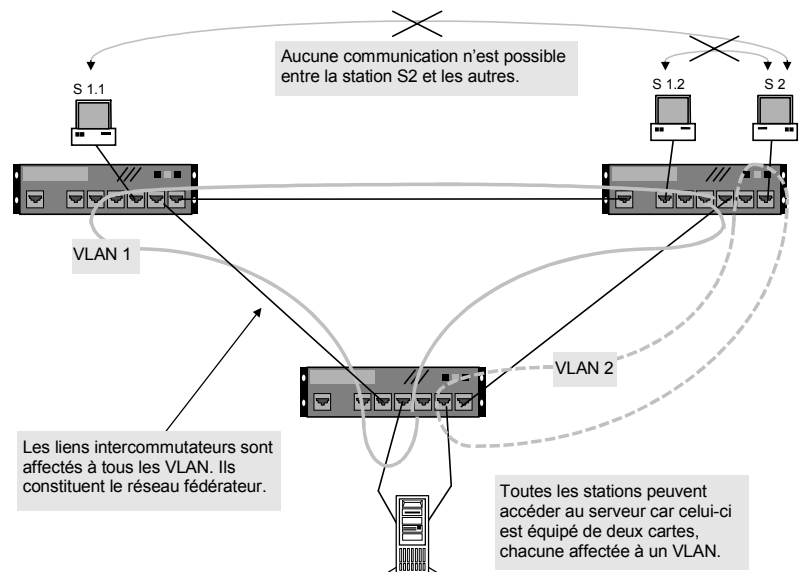
Mais, de nos jours, la traditionnelle répartition 80/20 (80 % du trafic local sur le réseau et 20 % vers d'autres réseaux) n'est plus valable. Par exemple, la constitution de groupes de travail pluridisciplinaires amène des personnes dispersées au sein de l'immeuble à établir des liens de communication privilégiés entre elles.

Traduit en termes techniques, les flux réseau générés de manière privilégiée entre les postes de travail et les serveurs évolueront sans cesse.

En outre, la constitution de réseaux isolés regroupant des utilisateurs géographiquement dispersés pourrait s'avérer nécessaire.

La solution à ces besoins passe par les **VLAN** (*Virtual Local Area Network*). Cette technologie permet de définir des segments Ethernet logiques, indépendamment de la localisation géographique des postes de travail. Une trame émise au sein d'un VLAN ne sera diffusée qu'aux stations participant audit VLAN.

Figure 11-1.
Principe des VLAN.



Or, seuls les commutateurs permettent de créer des VLAN, c'est-à-dire de segmenter le réseau correspondant à des groupes d'utilisateurs indépendamment de leur localisation géographique.

De plus, nous avons vu au chapitre 6 que les commutateurs sont nécessaires pour des applications multimédias (téléphonie et visioconférence sur IP), pour de gros volumes de données et pour une question de fiabilité.

Si l'on veut répondre à tous ces besoins, il est donc nécessaire d'installer des commutateurs sur l'ensemble de notre réseau (eau et gaz à tous les étages, pour ainsi dire !). Le choix de ces équipements s'impose donc à la fois pour des questions de performances et d'architecture.

Nous choisissons donc la solution 100 % commutateurs.

Routeur ou commutateur de niveau 3 ?

L'intérêt de partitionner notre réseau en réseaux plus petits est de circonscrire localement les flux générés par un groupe d'utilisateurs, de créer des zones isolées, ou encore de réduire les flux générés par les broadcast (surtout pour les grands réseaux).

La constitution de réseaux distincts (constituant chacun un domaine de broadcast MAC) nécessite cependant de les interconnecter quelque part. En effet, même s'ils appartiennent à des groupes différents, les utilisateurs doivent, à un moment ou à un autre, accéder à des ressources communes (serveur d'annuaire, passerelles fax, base de données centrale, PABX, accès Internet, etc.).

Une fonction de **roulage** est donc nécessaire pour interconnecter ces différents réseaux, qu'ils soient physiquement ou virtuellement constitués.

À ce niveau, nous avons le choix entre deux types d'équipements : les routeurs et les commutateurs de niveau 3.

Comparé au commutateur de niveau 3, le routeur présente un certain nombre de désavantages : il est nettement moins performant et, de ce fait, dispose rarement d'interfaces Gigabit. Par ailleurs, il ne sait pas gérer les VLAN.

Ce dernier point mentionné implique que le routeur dispose d'autant d'interfaces qu'il y a de VLAN (si ceux-ci sont créés par port), ou d'autant d'adresses IP sur une interface qu'il y a de VLAN créés par adresse IP.

Pour les petits et moyens réseaux (moins de 800 postes) sans liens gigabit, on peut envisager un routeur pour interconnecter quelques VLAN. Au-delà de ces restrictions, le commutateur de niveau 3 s'impose.

LES COMMUTATEURS DE NIVEAUX 2 ET 3

Un commutateur de niveau 2 agit au niveau des couches physique et logique (niveaux 1 et 2). Il ne traite que les trames MAC. On parle de commutation de niveau 2 ou layer 2 switching.

Un commutateur de niveau 3, quant à lui, agit au niveau de la couche réseau (niveau 3). Il ne traite que les paquets IP. C'est l'équivalent d'un routeur mais en beaucoup plus performant. On parle de commutation de niveau 3 ou layer 3 switching.

Quelle architecture ?

Nous voilà donc confortés dans le choix des commutateurs. Mais quelle architecture retenir ? Et à quel débit ?

On le voit, pour notre réseau de 1500 postes, de nouvelles considérations viennent compliquer notre tâche, de nouveaux paramètres influent sur le choix de l'architecture. En fait, tout tourne autour du fédérateur, pièce maîtresse du réseau. Résumons :

1. Une architecture basée uniquement sur des commutateurs de niveau 2 a le mérite de la simplicité. Elle a été étudiée au chapitre 6. Si l'on veut créer des réseaux séparés, il faut employer des VLAN par port ou par adresses MAC, ce qui augmente la complexité d'exploitation.
2. Une architecture basée uniquement sur des commutateurs de niveau 3 est plus coûteuse. Elle est cependant plus souple que la précédente, car on peut choisir les classes d'adresses IP et les combiner. L'architecture est identique à celle de la première solution, seule la technologie change.
3. Une architecture basée sur des routeurs est la moins performante de toutes et la moins souple (pas de VLAN possible). En fait, les routeurs sont plutôt destinés aux réseaux WAN.
4. Une architecture reposant sur un réseau fédérateur ATM est la plus complexe et la plus fragile, car elle impose une combinaison de plusieurs technologies. Son débit est, de plus, limité à 622 Mbit/s, ce qui est un handicap certain face au Gigabit Ethernet.

En fait, le routage n'est nécessaire qu'au niveau du réseau fédérateur, car tous les commutateurs d'étage y seront reliés.

En définitive, le choix se portera sur des commutateurs de niveau 2 pour les étages, et des commutateurs de niveau 3 pour le réseau fédérateur. Dans la pratique, ces derniers sont également des commutateurs de niveau 2 équipés de cartes de commutation de niveau 3.

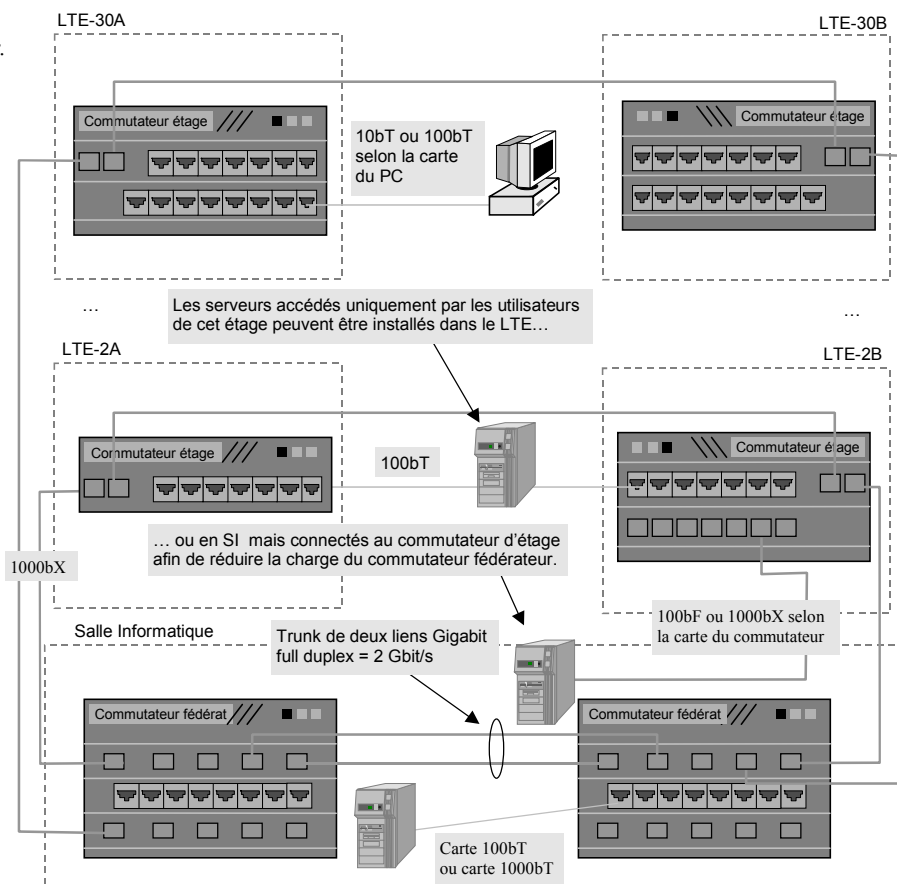
Pour le reste, nous appliquerons les recettes indiquées au chapitre 6.

Les commutateurs d'étage sont équipés de cartes 10/100bT ainsi que de deux ports uplink gigabit. Ils peuvent éventuellement être dotés de cartes 100bF ou de cartes gigabit pour connecter des serveurs délocalisés.

Les commutateurs fédérateurs sont principalement équipés de cartes gigabit pour être raccordés, d'une part, entre eux et, d'autre part, aux commutateurs d'étage. Ils peuvent éventuellement être dotés de cartes 10/100bT ou 1000bT, afin de connecter des serveurs situés dans des salles informatiques. Les cartes 1000bT offrent, en effet, une plus grande densité de port que leurs équivalents en fibre optique.

Les cartes en fibre optique sont utilisées partout où les distances sont supérieures à 90 mètres. Leur emploi est cependant systématisé au niveau du réseau fédérateur, même en dessous de cette distance, afin de disposer de configurations homogènes.

Figure 11-2.
Réseau fédérateur.



Configurer les VLAN

Même si les commutateurs fédérateurs sont équipés de cartes de commutation niveau 3, tous assurent la commutation de niveau 2. Le spanning tree doit donc être configuré sur tous les commutateurs, comme indiqué au chapitre 7.

De la même manière, un VLAN doit être configuré sur tous les commutateurs, afin qu'il soit connu de tous. Sur nos équipements (de marque Cisco), la création d'un VLAN par port s'effectue de la façon suivante :

```
set vlan 100 name VLAN_principal
set vlan 100 2/1-48
```

Les ports 1 à 48 situés sur la carte n° 2 seront ainsi affectés au VLAN 100 que nous avons appelé "VLAN principal".

L'opération suivante consiste à activer le protocole **802.1q** entre tous nos commutateurs, afin d'étendre la portée du VLAN à l'ensemble de notre réseau. Ce protocole ne doit être activé

que sur les ports qui raccordent des commutateurs entre eux, qu'on appellera des ports **trunk** (ports de liaison) :

```
set trunk 3/1 dot1q
set trunk 3/2 dot1q
```

Active le protocole 802.1q
sur ces ports

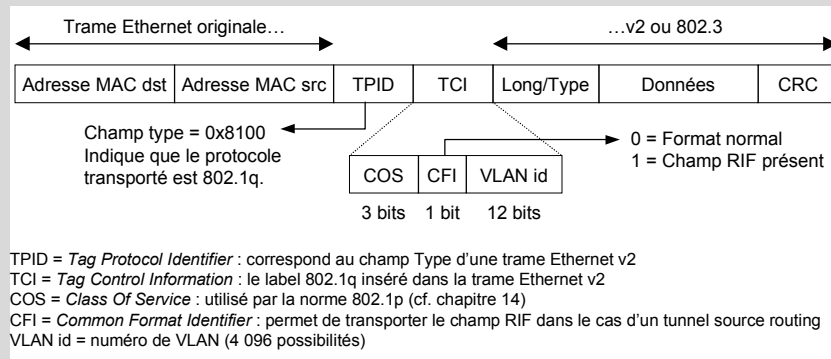
Chaque VLAN correspond à un domaine de broadcast et fonctionne donc avec un spanning tree indépendant. Il est donc possible de configurer ce protocole indépendamment pour chaque VLAN (voir chapitre 7) :

```
set spantree enable 100
set spantree fwddelay 15 100
set spantree hello 2 100
set spantree priority 16384 100
```

N° de VLAN

LE POINT SUR LES VLAN (IEEE 802.1q)

La norme 802.1q consiste à ajouter un champ à l'en-tête de la trame Ethernet initiale (802.3) à la fois pour gérer les VLAN et pour gérer des classes de service (802.1p).



Cette trame est véhiculée uniquement entre les commutateurs. Un VLAN peut donc être étendu à tout un réseau de commutateurs. Ces derniers ôtent le champ 802.1q lorsqu'ils transmettent la trame à un équipement terminal (PC, serveur, etc.) de manière que ces derniers retrouvent une trame conforme à la norme 802.3 ou Ethernet v2. La constitution des VLAN dépend de l'implémentation qui en est faite au sein des commutateurs. Il est ainsi possible de créer des VLAN :

- **par port** : toute trame entrant par un port est affectée d'office à un VLAN ;
- **par adresse MAC source** : toute trame disposant d'une telle adresse est affectée à un VLAN ;
- **par protocole** : toute trame véhiculant de l'IP, par exemple, est affectée à un VLAN ;
- **par adresse IP source** : toute trame véhiculant un paquet IP avec une telle adresse est affectée à un VLAN.

Un processus **spanning tree** (802.1d) est créé par VLAN. Par conséquent, les trames de **broadcast** et de **multi-cast** MAC émises au sein d'un VLAN ne seront pas propagées aux autres VLAN. En outre, les stations d'un VLAN ne pourront pas communiquer avec celles appartenant à un autre VLAN. Pour permettre cette fonction, il faut interconnecter les VLAN à l'aide d'un **routeur** ou d'un **commutateur de niveau 3**.

Les ports gigabit peuvent, de plus, être configurés de manière à opérer un contrôle de flux. Cela consiste en un signal envoyé à un autre commutateur pour lui demander de ralentir temporairement l'envoi de trames :

```
set port flowcontrol send 0/0-1 on
set port flowcontrol receive 0/0-1 on
```

Envoie...

...et accepte les signaux de contrôle de flux.

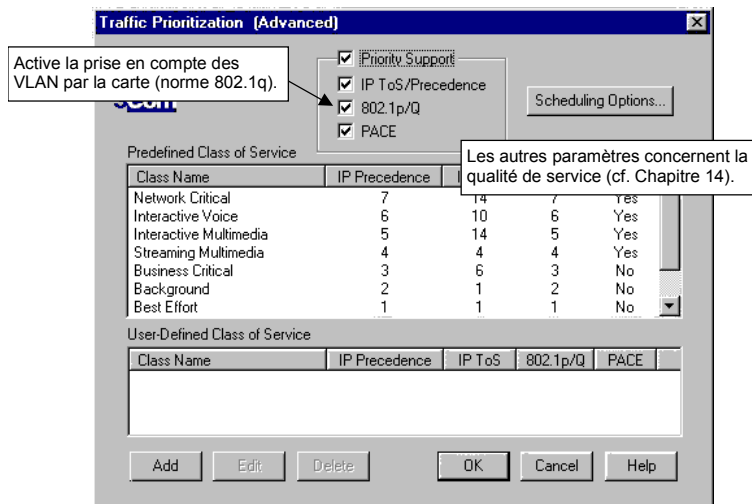
Dans notre architecture, il est également prévu d'agrégier deux liens gigabit entre les deux commutateurs fédérateurs :

```
set vlan 100 2/0-1
set port channel 2/0-1 desirable
set trunk 2/0 desirable dot1q
```

Crée un groupe de 2 ports

L'activation du trunking 802.1q sur un port est appliquée à tout le groupe.

La configuration des VLAN au niveau des cartes Ethernet des PC et des serveurs est possible si elles supportent le protocole 802.1q, ce qui est le cas de nos cartes 3com.



Il est cependant préférable de ne pas utiliser cette facilité pour les raisons suivantes :

- cela ajouterait une complexité supplémentaire aux tâches d'administration : il faudrait configurer à distance toutes les cartes des PC ;
- les utilisateurs trouveraient toujours le moyen de modifier la configuration de leur carte de manière à changer de VLAN ;
- pour éviter cela, il faudrait réaliser un contrôle au niveau des commutateurs, ce qui induirait une double exploitation.

Nous pourrions également créer des VLAN dynamiques par adresse IP. Là encore, l'exploitation est délicate et les modifications de la part des utilisateurs sont toujours possibles. L'affectation des VLAN par port a le mérite d'être simple, de faire partie de la configuration normale des commutateurs et de maîtriser l'étendue du VLAN sur notre réseau.

Extension du réseau fédérateur

Nos deux commutateurs fédérateurs étaient suffisants pour accueillir les 15 réseaux d'étage. Maintenant, nous devons connecter un autre site situé à quelques centaines de mètres, puis deux autres situés à quelques kilomètres.

Si tous les bâtiments sont situés sur un terrain privé (par exemple, un campus universitaire), nous pouvons poser de la fibre optique comme nous l'entendons. Dans le cas contraire, soit un opérateur nous loue des câbles en fibre optique, soit nous devons obtenir une dérogation pour en poser entre nos bâtiments.

Dans tous les cas, nous supposons donc que les bâtiments sont reliés entre eux par des câbles en fibre optique. Car l'enjeu est maintenant d'étendre notre réseau fédérateur pour en faire un réseau de campus.

On parle également de MAN (*Metropolitan Area Network*), bien qu'aucune technologie particulière, autre que celle utilisée en LAN, ne soit associée à ce type de réseau. Il s'agit simplement d'une dénomination conceptuelle.

Pour ceux qui en doutaient encore, le Gigabit Ethernet convient parfaitement à ce type de besoins. Certains opérateurs proposent même ce service sur quelques centaines de kilomètres. Tout dépend de la fibre optique utilisée.

	Support de transmission	Distance
1000bSX	Fibre multimode 62,5µ	De 2 à 300 m
	Fibre multimode 50µ	De 2 à 550 m
1000bLX	Fibre multimode 62,5µ	De 2 à 550 m
	Fibre multimode 50µ	De 2 à 550 m
	Fibre monomode 9µ	De 2 à 3 000 m

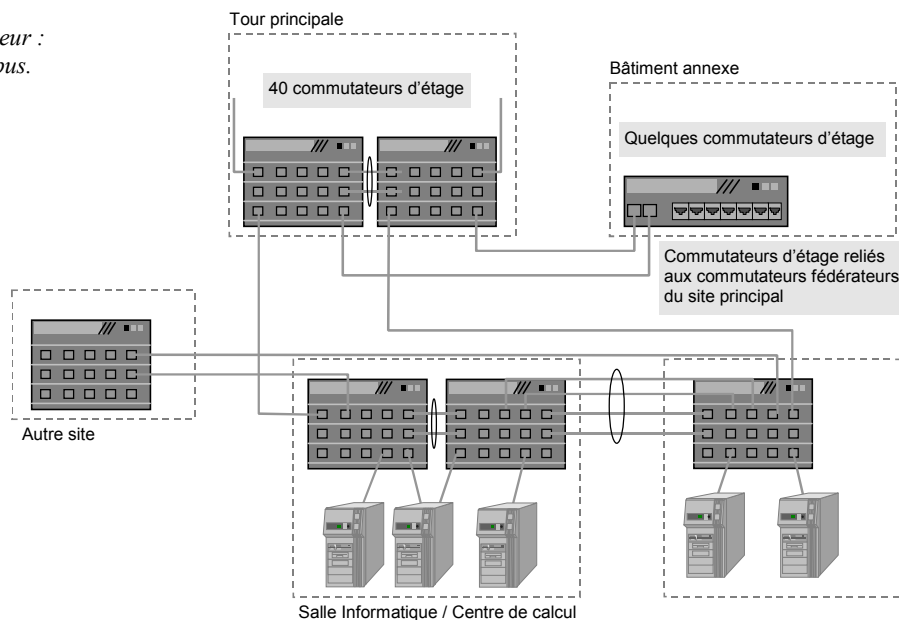
Côté performances, le Gigabit Ethernet est à la hauteur des débits annoncés :

- un débit réel de 761 Mbit/s pour des trames de 64 octets (soit 1 488 095 paquets par seconde) ;
- un débit réel de 986 Mbit/s pour des trames de 1 518 octets (soit 81 274 paquets par seconde).

Confortés dans notre choix du gigabit, nous nous retrouvons avec plusieurs commutateurs fédérateurs à interconnecter.

Figure 11-3.

*Extension
du réseau fédérateur :
le réseau de campus.*



Nous aurions pu mailler tous les commutateurs fédérateurs de manière à offrir des routes multiples. Cela est envisageable si les serveurs sont disséminés dans différents bâtiments.

Quand cela est possible, il est cependant préférable de respecter les principes suivants :

- Choisir deux commutateurs fédérateurs de campus qui fédéreront également les autres commutateurs fédérateurs de site (ou en dédier deux autres), de manière à centraliser les flux intersites au sein d'un nombre réduit de matrices de commutation. Ces deux équipements peuvent être situés dans deux bâtiments différents.
- Relier les deux commutateurs fédérateurs de campus par un lien très haut débit, dans notre cas quatre liens gigabits.
- Connecter les fédérateurs de site aux fédérateurs de campus par deux liens distincts en partage de charge et en redondance, de préférence sur deux commutateurs distincts, de manière à pallier la défaillance d'un équipement.

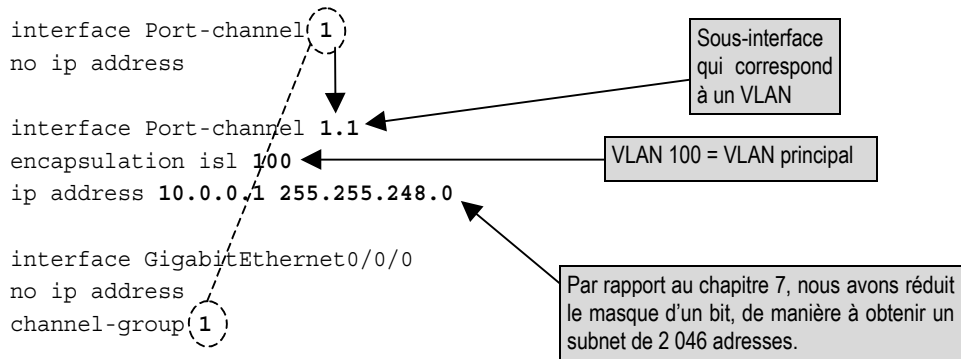
L'extension du réseau fédérateur se fait donc de manière très simple, sans remettre en cause les choix technologiques et l'architecture.

Il est à noter que, grâce aux VLAN, cette architecture permet à plusieurs sociétés de cohabiter sur la même infrastructure tout en étant isolées.

L'adressage et le routage IP

Nous avons décidé de créer un VLAN et d'affecter de manière statique les ports à ce VLAN, que nous avons appelé VLAN principal. D'autres VLAN peuvent être créés pour des réseaux dédiés.

Cela implique d'affecter un subnet IP à chaque VLAN, et par conséquent de configurer nos cartes de commutation niveau 3 (fonctionnellement équivalentes, rappelons-le, à des routeurs). Pour notre VLAN principal, cela est réalisé comme suit, conformément à notre plan d'adressage établi au chapitre 7 :



L'encapsulation "isl" (*Inter-Switch Link*) active le protocole propriétaire Cisco équivalent de la norme 802.1q. Dans ce cas particulier, nous ne pouvons faire autrement.

La carte de commutation de niveau 3 dispose ainsi d'un attachement sur le VLAN principal, qui est réalisé au niveau de la matrice de commutation. La création de tout autre VLAN sera réalisée sur le même modèle.

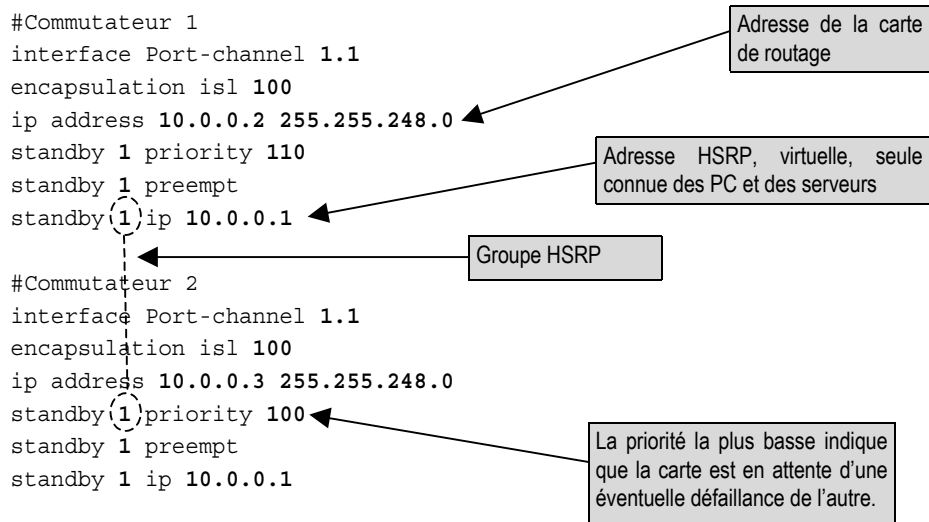
Le routage entre VLAN, et donc entre subnets IP, est effectué par la carte de commutation. Pour sortir du VLAN principal, les PC et les serveurs doivent connaître la route de sortie, c'est-à-dire la route par défaut (*default gateway*), comme cela était le cas au chapitre 8. Mais, cette fois, elle doit pointer sur l'adresse IP de la carte de commutation, à savoir 10.0.0.1.

Sur chaque VLAN, la passerelle par défaut des PC et des serveurs pointe donc sur l'adresse IP du commutateur de niveau 3 attaché audit VLAN.

Redondance du routage

Si, comme sur notre site parisien, nous disposons de deux commutateurs fédérateurs, chacun équipé d'une carte de commutation de niveau 3 (carte de routage), il est intéressant d'assurer la redondance de la route par défaut vis-à-vis des PC et des serveurs. Sur nos équipements, cela est réalisé grâce à la fonction HSRP (*Hot Standby Router Protocol*).

Le principe repose sur un groupe de n routeurs (ou cartes de commutation de niveau 3 dans notre cas), dont l'un est désigné actif. Comme d'habitude, chaque interface est associée à une adresse IP et à une adresse MAC. Mais le routeur actif reçoit en plus une adresse IP (définie comme route par défaut) associée à une adresse MAC qui seule répond au protocole de résolution d'adresses ARP (voir chapitre 7). En cas de défaillance du routeur actif, un nouveau routeur est élu parmi les $N-1$ restants en fonction des priorités affectées au sein du groupe HSRP, qui s'approprie les adresses HSRP (MAC et IP) :



La route par défaut configurée sur les PC et serveurs est celle de l'adresse HSRP, à savoir 10.0.0.1.

Ce principe peut être appliqué à chaque VLAN. Il est alors conseillé d'affecter les priorités de telle manière que chacune des cartes de commutation de niveau 3 soit active au moins pour un VLAN, et ce afin de répartir la charge de routage.



Certaines piles IP, comme celle de Windows NT, intègrent un mécanisme de détection de panne du routeur par défaut (*dead gateway detection*), tel que décrit dans le RFC 816. Si la station constate qu'elle ne parvient plus à joindre son routeur par défaut, elle en choisira un autre parmi une liste définie dans le menu des propriétés de TCP/IP, case "Avancé...", section "Passerelle". Pour activer ce mécanisme de détection, il faut positionner à "1" la clé de registre "HKEY_LOCAL_MACHINE\System\CurrentControlSet\Services\Tcpip\Parameters\EnableDeadGWDetect".

LE POINT SUR VRRP (RFC 2338)

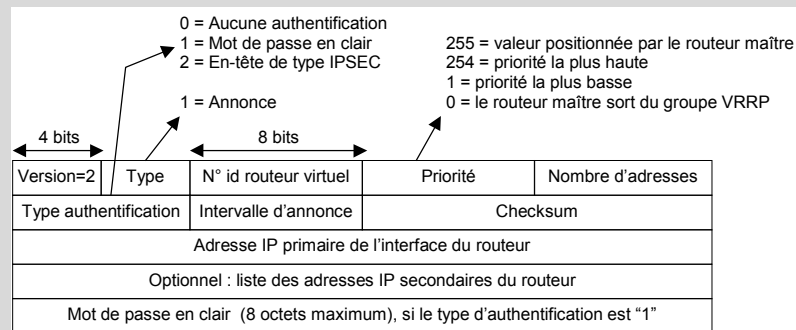
Le protocole **VRRP** (*Virtual Router Redundancy Protocol*) reprend les principes du protocole **HSRP** (*Hot Standby Router Protocol* – RFC 2281) spécifié par Cisco. Les formats des paquets sont en revanche différents, ce qui rend ces deux protocoles incompatibles.

Le but est ici d'offrir une **redondance** de routeurs pour les machines (notamment les PC) utilisant le mécanisme de **passerelle par défaut** (*default gateway*). Même si plusieurs routeurs sont connectés à un segment Ethernet, la passerelle par défaut des PC pointera vers une seule adresse IP, celle d'un des routeurs choisis par l'administrateur du réseau.

Le protocole VRRP permet d'ajouter, en plus des adresses propres à chaque routeur, une **adresse IP virtuelle** vers laquelle les passerelles par défaut peuvent pointer. À un instant donné, seul le routeur désigné **maître** détiendra l'adresse virtuelle, et pourra assurer le traitement des paquets à destination de cette adresse.

Ainsi, lorsque la pile IP du PC devra résoudre, grâce à ARP, l'adresse de sa passerelle par défaut, seul le routeur maître répondra en indiquant l'**adresse MAC virtuelle**.

Le routeur maître envoie un paquet d'annonce à intervalle régulier. Si les autres routeurs n'en reçoivent plus au bout de l'intervalle de temps spécifié dans le dernier paquet reçu (par défaut une seconde), ils considèrent que le routeur maître est en panne et entrent alors dans un processus d'**élection** en envoyant des annonces. Celui dont la **priorité** est la plus élevée, devient alors maître et prend le contrôle de l'adresse virtuelle.



Les paquets VRRP disposent du numéro de **protocole 112**. Ils sont envoyés dans des paquets IP à destination de l'adresse multicast **224.0.0.18**, dont l'adresse source est la véritable adresse IP du routeur et dont le TTL est obligatoirement fixé à **255**. Le tout est envoyé dans une trame MAC d'adresse source **00-00-5E-00-01-xx** où "xx" représente le numéro d'identification du routeur virtuel (identique au champ "n° id routeur virtuel" du paquet VRRP).

De son côté, HSRP fonctionne au-dessus d'UDP avec l'adresse multicast 224.0.0.2 et un TTL fixé à 1. L'adresse MAC virtuelle utilisée est 00-00-0C-07-AC-xx.

Un même routeur peut participer à plusieurs groupes VRRP et plusieurs groupes VRRP peuvent cohabiter sur un LAN. Si, via le système des priorités, on s'arrange pour que chaque routeur d'un LAN soit maître pour un groupe, et si on répartit les passerelles par défaut des PC sur chacune des adresses virtuelles, il est alors possible de **partager la charge** de routage entre les routeurs.

Les mécanismes classiques d'icmp-redirect et de proxy ARP sont toujours opérants.

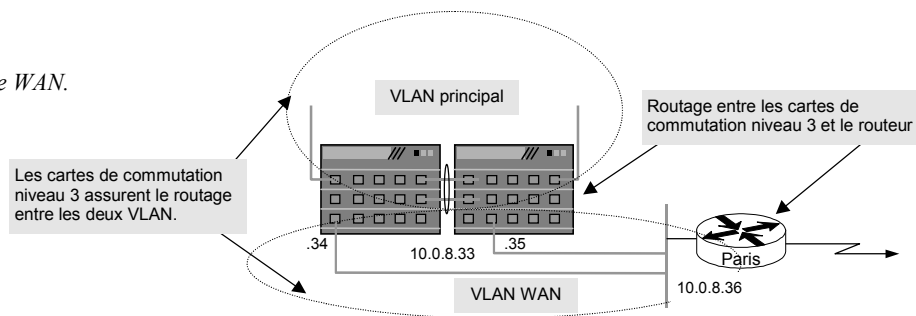
La rencontre du LAN et du WAN

Les routeurs qui interconnectent notre site aux autres peuvent directement être raccordés au réseau fédérateur. Nous préférons cependant créer un VLAN spécifique, afin de marquer la frontière entre les deux mondes, ce qui procure certains avantages :

- tout changement d'architecture ou de configuration du LAN n'affecte pas le WAN, et inversement ;
- la gestion du WAN peut être centralisée à partir d'un autre site, ce dernier gardant son autonomie sur le LAN ;
- idem lorsque les routeurs sont gérés par un ou plusieurs opérateurs.

Figure 11-4.

La frontière entre le LAN et le WAN.



Comme pour le VLAN principal, il faut affecter un subnet IP au VLAN WAN, puis des adresses IP aux cartes de commutation et à HSRP. Un subnet de 30 adresses, pris dans notre plan d'adressage, sera largement suffisant :

```
interface Port-channel 1.2
encapsulation isl 5
ip address 10.0.9.34 255.255.255.224
```

VLAN n° 5 = VLAN WAN

```
standby 2 priority 110
standby 2 preempt
standby 2 ip 10.0.9.33
```

Groupe HSRP n° 2 affecté à ce VLAN

Il suffit ensuite de configurer le routage entre nos commutateurs LAN et le routeur WAN. La manière la plus simple de le faire est de définir des routes statiques, soit une par défaut, soit explicitement pour chaque site distant connu :

```
# route par défaut
ip route 0.0.0.0 0.0.0.0 10.0.9.36
# OU routes statiques explicites
ip route 10.4.0.0 255.252.0.0 10.0.9.36
ip route 10.8.0.0 255.252.0.0 10.0.9.36
```

Tout ce qui n'est pas connu est
envoyé au routeur

Vers Toulouse et vers Strasbourg
via le routeur

Inversement, nous indiquons au routeur comment joindre le VLAN principal du site parisien :

```
ip route 10.0.0.0 255.252.0.0 10.0.9.33

int e0
ip address 10.0.9.36 255.255.255.224
```

Sur Toulouse, nous avons deux routeurs WAN qui peuvent être redondants pour des liaisons Frame-Relay. Il est alors possible de configurer HSRP à la fois sur les routeurs WAN et sur les routeurs LAN (les cartes de commutation de niveau 3).

Le routage sur le WAN

Une fois arrivés sur le WAN, les paquets IP se trouvent face à de multiples routes allant vers la même destination.

Il est envisageable de programmer tous les routeurs avec des routes statiques, comme nous l'avons fait précédemment, mais cette tâche peut s'avérer complexe et fastidieuse, surtout s'il faut envisager des routes de secours.

Sur le WAN, le plus simple est d'utiliser un protocole de **routage dynamique**. Nous avons alors le choix entre RIP et OSPF (voir encadré). Ce dernier est cependant le plus performant et le plus répandu, même s'il est un peu plus complexe à programmer. Nous choisissons donc OSPF.

Configuration du routage

La première tâche est d'activer le routage OSPF. Sur nos routeurs Cisco, il faut attribuer un numéro de processus, car plusieurs instances d'OSPF peuvent fonctionner simultanément :

```
router ospf 1
```

Avec OSPF, la première tâche est de définir l'aire 0, appelée *backbone area*.

QU'EST-CE QU'UN PROTOCOLE DE ROUTAGE ?

Le routage est l'action de commuter les paquets d'un réseau IP à l'autre en fonction de leur adresse IP de destination. Le routeur se base sur des routes **statiques** (configurées par l'administrateur) et **dynamiques** (appries par des protocoles de routage).

Le routeur maintient ainsi une base de données des coûts des routes associées, ce qui permet de calculer le meilleur chemin.

Afin de réduire le trafic réseau généré par les protocoles de routage, de réduire la taille des bases de données et de déléguer l'administration, les réseaux IP sont découpés en domaines appelés **systèmes autonomes** (AS, *Autonomous System*).

Les protocoles spécialisés dans le routage au sein d'un AS sont de type **IGP** (*Interior Gateway Protocol*). Les plus courants sont **RIP** (*Routing Information Protocol*) et **OSPF** (*Open Shortest Path First*). Les protocoles spécialisés dans le routage inter AS sont de type **EGP** (*Exterior Gateway Protocol*). Le plus répandu est **BGP** (*Border Gateway protocol*).

Au sein d'un AS, tous les routeurs disposent de la même base de données décrivant la topologie de l'AS.

Les IGP utilisent deux types d'algorithmes pour calculer les routes : celui à **vecteur de distance** (Bellman-Ford) utilisé par RIP, et celui de l'**arbre du plus court chemin** (Dijkstra), plus performant, qui est utilisé par OSPF.

Même si de multiples configurations sont possibles avec OSPF, il est cependant conseillé de respecter les règles suivantes :

- L'aire 0 doit couvrir toutes les interfaces WAN des routeurs (c'est-à-dire les interfaces série, Frame-Relay, ATM, LS, RNIS, etc.).
- Une aire doit être définie par site ou par groupe de sites fédérés autour d'un campus. L'intérêt est de pouvoir contrôler la diffusion des routes, par exemple, d'empêcher qu'un subnet parisien puisse être vu des autres sites.

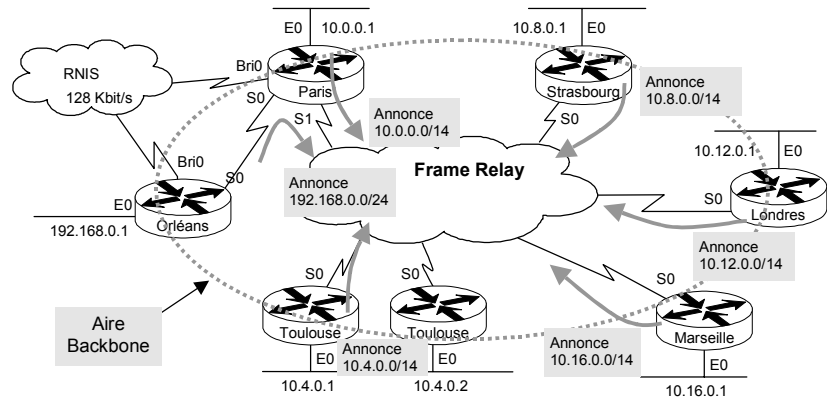
Étant donné que notre plan d'adressage défini au chapitre 7 prévoit l'affectation d'un subnet complet à l'ensemble des liaisons WAN, une seule commande sur chaque routeur est nécessaire pour affecter l'aire 0 :

```
network 172.16.0.0 0.0.255.255 area 0.0.0.0
```

Tous les réseaux WAN sont
dans l'aire backbone

L'aire OSPF est un numéro sur 32 bits qui peut être noté à la manière d'une adresse IP. La notation du masque associé au subnet à annoncer utilise, quant à elle, une convention inverse à celle utilisée pour les adresses IP (les bits à "0" indiquent la partie réseau).

Figure 11-5.
Configuration OSPF.



Côté LAN, il n'y a pas de contrainte particulière à l'affectation d'une aire. Nous choisissons d'en affecter une par site (ou par campus) si cela se révélait nécessaire.

Aire OSPF	Site
0.0.0.1	Région parisienne
0.0.0.2	Région toulousaine
0.0.0.3	Strasbourg
Etc.	...

C'est justement le cas à Toulouse, car nous avons deux routeurs, connectés à l'aire 0 d'un côté et au même réseau local de l'autre. Afin que ces deux routeurs puissent échanger leurs tables de routage et se secourir mutuellement, il faut positionner leur interface locale dans une aire.

S'il n'y a que deux routeurs, le plus simple est de tout mettre dans l'aire 0. Si le réseau de Toulouse grandit au point d'intégrer plusieurs routeurs (ou cartes de commutation de niveau 3), on peut envisager de créer une aire sur ce site, afin de réduire le trafic sur le WAN et de mieux contrôler la diffusion des routes :

```
network 10.4.0.0 0.3.255.255 area 0.0.0.2
```

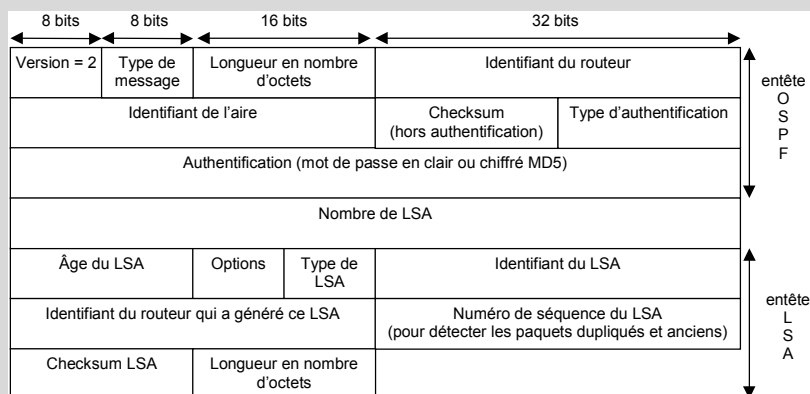
Redondance en cas de panne

En reprenant le réseau Frame-Relay que nous avons construit au chapitre 10, nous pouvons apercevoir que, en cas de panne du routeur de Strasbourg, l'aire backbone serait coupée en deux, empêchant toute diffusion des routes. Même s'il y a continuité du réseau local, l'aire backbone est séparée par l'aire 2.

LE POINT SUR OSPF (RFC 2328)

Le protocole OSPF (*Open Shortest Path First*) découpe l'AS (*Autonomous System*) en **aires**. Toutes les aires doivent être adjacentes à l'aire 0 (**backbone area**) qui doit être contiguë. Si elle ne l'est pas, un **lien virtuel** doit être configuré pour assurer sa continuité logique. Les paquets routés entre aires doivent tous passer par la backbone area *via* les **routeurs de bordure**.

Les routeurs diffusent régulièrement des messages d'annonce **LSA** (*Link State Advertisement*) pour indiquer quels réseaux leur sont directement attachés. Les LSA sont diffusés à tous les routeurs de l'aire ; ils permettent à chacun d'entre eux de disposer de la même base de données d'**état des liens** et de calculer l'**arbre du plus court chemin** dont il est la racine. Un routeur gère autant de bases de données et calcule autant d'arbres qu'il y a d'aires auxquelles il est connecté.



Les routeurs diffusent régulièrement des messages **Hello** afin d'annoncer leur présence à leurs voisins sur les réseaux multipoints supportant le broadcast (par exemple Ethernet). Celui dont la priorité est la plus grande est élu **routeur désigné** ; il a la charge d'inclure ce réseau dans ses LSA. Sur les réseaux Ethernet, les messages sont envoyés dans des paquets multicast 224.0.0.5.

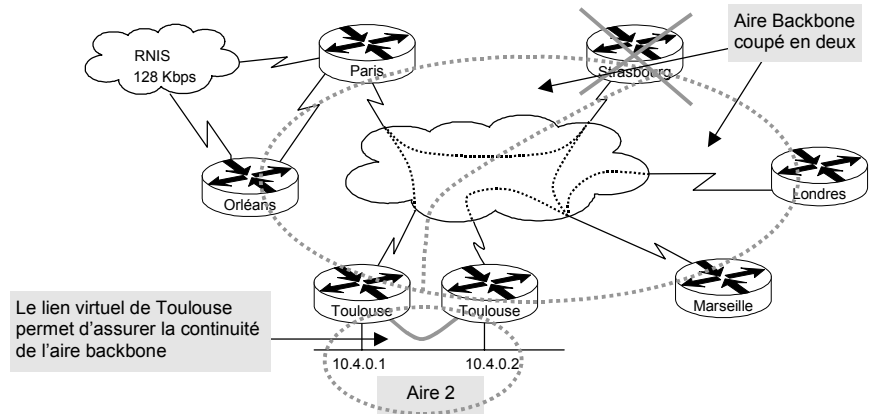
Dans la **backbone area**, les routeurs de bordure s'échangent les bases de données des aires auxquelles ils sont rattachés. Ils calculent les meilleures routes qui sont ensuite diffusées aux routeurs intra-aire. Les routeurs intra-aire calculent la meilleure route pour sortir de l'aire (*via* un routeur de bordure). Le **routeur frontière** (qui peut être situé n'importe où dans l'AS) assure le même rôle pour les routes permettant de sortir de l'AS.

Quatre types de LSA sont échangés :

- **router-LSA** : émis par tous les routeurs d'une aire pour décrire l'état et indiquer le coût de leur interface.
- **network-LSA** : émis par les routeurs désignés pour annoncer les réseaux de type broadcast (Ethernet, par exemple) ;
- **summary-LSA** : émis par les routeurs de bordure ;
- **AS-external-LSA** : émis par les routeurs de frontière.

Aucun AS-external-LSA n'est envoyé dans les aires configurées en **stub area**. À la place, le routeur de bordure diffuse une route par défaut.

Figure 11-6.
Liens virtuels OSPF.



Il est à noter que ce cas de figure n'existerait pas si Toulouse ne disposait que d'un seul routeur ou si les interfaces Ethernet des deux routeurs étaient situées dans l'aire 0.

La solution à ce problème passe par la création d'un lien virtuel entre les deux routeurs de Toulouse :

```
# Routeur 1
```

```
area 2 virtual-link 10.4.0.2
```

```
# Routeur 2
```

```
area 2 virtual-link 10.4.0.1
```

Aire de transit commune
aux deux routeurs

Adresse du routeur à l'autre
bout du lien virtuel

Ce lien permet d'assurer la continuité de l'aire backbone *via* l'aire de transit de Toulouse.

Ajustement des paramètres

Diffuser les routes statiques

Certains sites peuvent comporter des routeurs configurés uniquement avec des routes statiques. Si ces routes doivent être connues des autres sites, il est alors impératif de les diffuser au processus OSPF, de manière que ce dernier les diffuse dynamiquement à ses voisins :

```
router ospf 1
```

```
redistribute static
```

Modifier le coût des routes

Pour calculer le coût des routes, et donc choisir la meilleure, OSPF se base sur la bande passante du lien. Sur nos routeurs, il est nécessaire de l'indiquer manuellement, par exemple 512 Kbit/s sur les routeurs de Toulouse :

```
int s0
bandwidth 512
```

Par défaut, le coût associé à l'interface est de 100 000 divisé par le débit exprimé en Kbit/s, ce qui donne, par exemple, un coût de 1 562 pour un débit de 64 Kbit/s. Il est néanmoins possible de le modifier, comme suit :

```
int s0
ip ospf cost 300
```

Valeur de 1 à 65 535

Limiter la diffusion des routes

Dans certains cas, il peut être intéressant de limiter la diffusion de certaines routes afin qu'elles ne soient pas connues d'autres sites, par exemple pour des questions de confidentialité ou pour forcer le chemin à emprunter :

```
router ospf 1
distribute-list 11 out
```

Le routeur parisien ne diffuse pas le réseau du site d'Orléans.

```
access-list 11 deny 192.168.0.0 0 0.0.0.255
access-list 11 permit any
```

De la même manière, un routeur peut ne pas accepter une route si, par exemple, le site d'Orléans doit être caché uniquement à celui de Londres :

```
router ospf 1
distribute-list 11 in
```

Le routeur de Londres filtre la route du site d'Orléans.

```
access-list 11 deny 192.168.0.0 0 0.0.0.255
access-list 11 permit any
```

Modifier la fréquence des échanges

Les routeurs OSPF voisins s'échangent des paquets Hello selon une périodicité qu'il est possible de modifier :

```
int s0
ip ospf hello-interval 10
ip ospf dead-interval 40
```

Envoi un paquet Hello à ses voisins toutes les 10 secondes.

Le routeur voisin est déclaré absent au bout de 40 secondes (par défaut, 4 x le hello-interval).

Forcer l'élection du routeur désigné

Lorsque, comme cela est le cas à Toulouse, il existe deux routeurs sur le même réseau Ethernet, seul le routeur désigné va diffuser le subnet IP de l'aire n° 2. Est élu "désigné" le routeur dont la priorité est la plus haute ; en cas de niveau de priorité identique, c'est celui dont l'adresse IP est la plus haute :

```
ip ospf priority 1
```

Valeur par défaut

Les performances d'OSPF

Le RFC 1245 fournit quelques statistiques relevées sur les routeurs de l'Internet :

- Chaque entrée de la base d'états de liens est mise à jour toutes les 30 minutes en moyenne.
- Selon les cas, l'arbre du plus court chemin est recalculé toutes les 13 à 50 minutes.
- En moyenne, un paquet d'annonce contient trois LSA.
- Pour 2 000 entrées dans une base de données OSPF, la bande passante consommée par l'émission des LSA représente moins de 0,5 Kbit/s.

Type d'annonce	Taille moyenne dans les paquets	Mémoire routeur
External LSA	36 octets	64 octets
Router et Network LSA	108 octets	192 octets
Summary LSA	36 octets	64 octets
En-tête OSPF	24 octets	--
En-tête IP	20 octets	

Le temps CPU pour calculer l'arbre du plus court chemin (algorithme de Dijkstra) est de l'ordre de $n \cdot \log(n)$ pour N routes et 200 routeurs, soit environ 15 millisecondes pour un processeur de 10 Mips. En découpant un système autonome en aires, la charge CPU est réduite, car il y a moins de routeurs à prendre en compte, le calcul SPF étant réalisé au sein d'une aire.

TROISIÈME PARTIE

Se préparer au multimédia

12

Les flux multimédias

L'objectif de ce court chapitre est de présenter les caractéristiques des flux multimédias (essentiellement la voix et la vidéo), afin de montrer comment leurs particularités influent sur un réseau de paquet tel que TCP/IP.

Vous découvrirez ainsi :

- comment sont transportés le son et l'image sous forme numérique ;
- ce qu'est un codec ;
- les problèmes posés par les délais de transit et la gigue ;
- quels débits ces types de flux engendrent sur votre réseau.

Les caractéristiques des flux multimédias

Initialement, un réseau IP, tel que l'Internet, était conçu pour véhiculer des données entre deux machines : transfert de fichiers, connexion web, messagerie, etc. Depuis, la voix et l'image ont fait leur apparition et ont étendu le champ d'utilisations du réseau : téléphonie, diffusion de films à la demande et conférences à plusieurs. Dans ce dernier domaine, on distingue l'audioconférence (voix uniquement), la visioconférence (voix + vidéo) et, d'une manière générale, la téléconférence (voix + vidéo + données).

Ces flux multimédias induisent un certain nombre de contraintes nouvelles :

- Les signaux audio et vidéo doivent être numérisés, ce qui veut dire convertir les signaux analogiques en bits numériques.
- La voix nécessite une bonne synchronisation entre l'émetteur et le récepteur.
- La vidéo engendre une augmentation du volume des données transférées.
- La téléconférence nécessite de diffuser un flux entre un émetteur et plusieurs récepteurs qui peuvent devenir, à leur tour, émetteurs.

La numérisation des signaux nécessite d'échantillonner la voix, de la quantifier, de la coder et, de plus en plus souvent, de la compresser :

- L'**échantillonnage** consiste à prélever des échantillons du signal à intervalles réguliers, à l'instar du cinéma qui utilise 24 images par seconde pour traduire le mouvement. L'amplitude des échantillons prélevés peut varier de façon illimitée, mais doit pouvoir être représentée par un nombre fini de valeurs binaires. La plupart du temps, l'échantillon porte sur 8 bits, alors que, pour la haute fidélité (requis, par exemple, pour un Compact Disc), l'échantillon porte sur 16 bits.
- La **quantification** fait correspondre une valeur à l'amplitude d'un échantillon par rapport à des valeurs-étalons appelées niveaux de quantification. La valeur sur un octet sera celle du niveau de quantification le plus proche.
- Le **codage** consiste à transmettre un flux d'informations binaires correspondant à l'échantillon représenté par un octet.
- De plus en plus souvent, le codage est associé à un algorithme de **compression**.

Par exemple, la voix génère des signaux à une fréquence oscillant entre 300 Hz et 3 300 Hz, valeur arrondie à 4 000 Hz par les équipements numériques. Un chercheur, appelé Shannon, a montré que la fréquence d'échantillonnage devait être égale au double de la fréquence du signal à numériser. Un signal analogique de 4 000 Hz nécessite donc 8 000 échantillons par seconde, qui sont représentés sur 8 bits, soit un débit de 64 Kbit/s pour coder la voix. Cette unité, appelée **DS0** (*Digital Signaling 0*), a longtemps été la référence et continue encore de l'être pour les liaisons d'accès E1/T1 proposées par les opérateurs, aussi bien pour la voix que pour les données.

Les quatre opérations constituant la numérisation d'un signal analogique (audio ou vidéo) sont réalisées par des processeurs spécialisés appelés **DSP** (*Digital Signaling Processing*). Les algorithmes qui définissent la manière de réaliser ces opérations sont appelés **Codec** (codeur/décodeur).

Les codec audio

Par exemple, le codec utilisé sur les liaisons E1/T1 ainsi que sur le RTC répond à la norme G.711. Le codage est de type **PCM** (*Pulse Code Modulation*) ; il n'utilise aucun algorithme de compression. En France, la norme est appelée **MIC** (Modulation par impulsions codées), d'où le nom des liaisons à 2 Mbit/s proposées par France Télécom (32 canaux DS0, dont 2 pour la signalisation).

Le codec G.711 existe en deux variantes de codage : A-law (Europe) et μ -law (Amérique du Nord et Japon).

Codec (standard ITU)	Algorithme de codage	Échantillonnage	Débit réseau nécessaire	Délai de traitement du codec	Durée de l'échantillon	Taille de l'échantillon
G.711	PCM (A-law / μ -law)	8 kHz	64 kbit/s	0,75 ms		
G.722	ADPCM	7 kHz	64 kbit/s	1 ms		1 octet
G.723	MP-MLQ et ACELP	8 kHz	5,3 et 6,3 kbit/s	7,5 ms	30 ms	20 et 24 octets
G.726	ADPCM	8 kHz	16, 24, 32 et 40 kbit/s	2 à 3,5 ms		
G.728	LD-CELP	3,1 kHz	16 kbit/s	3 à 5 ms	3,1 ms	5 octets
G.729	CS-ACELP	4 kHz	8 kbit/s	15 ms	10 ms	10 octets

Au délai de traitement du codec, il faut ajouter le temps mis pour remplir une trame. Par exemple, le codec G.723 génère une trame contenant 30 ms de voix. Elle comporte 240 échantillons compressés à 189 bits (24 octets) ou à 158 bits (20 octets), les deux premiers bits indiquant respectivement le type de codec et la taille de la trame.

Les significations des sigles désignant les différents algorithmes de codage qui viennent d'être cités sont les suivantes :

- ADPCM (*Adaptive Differential Pulse Code Modulation*) ;
- CELP (*Code Excited Linear Prediction*) ;
- LD-CELP (*Low-Delay Code-Excited Linear-Prediction*) ;
- CS-ACELP (*Conjugate-Structure Algebraic Code-Excited Linear-Prediction*) ;
- MP-MLQ (*MultiPulse Maximum Likelihood Quantization*).

La qualité d'un codec dépend essentiellement de l'échantillonnage et du taux de compression : plus ce dernier est élevé, plus la voix apparaît métallique à celui qui l'écoute.

Sans compression, le codec G.711 sert de référence pour la qualité de la voix sur le téléphone, référence qui est appelée *toll quality*. La perception de la qualité de la voix est subjective. C’est pour cela qu’elle est mesurée à travers un indice moyen de satisfaction appelé MOS (*Mean Opinion Score*), sur une échelle allant de 0 (très mauvais) à 5 (très bon).

Codec	MOS	Application
G.711	4,1	RTC, RNIS
G.722		
G.723	3,65	IP
G.726	3,85	
G.728	3,81	
G.729	3,92	Frame-Relay

Les voix ayant une sonorité métallique obtiennent un MOS inférieur à 3 ; celles acceptables pour une communication téléphonique, un MOS compris entre 3 et 4. Au-delà de 4, la qualité est excellente.

Le choix d’un codec résulte donc d’un compromis qualité/débit sur le réseau utilisé. Ainsi, le programme Netmeeting permet-il de sélectionner le plus approprié à votre contexte (menu “Outil→Options→Audio→Avancé”).



Les codec vidéo

Une image vidéo affichée sur un écran de PC représente 640 × 480 points en 16,7 millions de couleurs (24 bits) × 25 images par seconde, ce qui correspond à un débit de 23 Mbit/s (640 × 480 × 3 × 25).

Afin de diminuer le débit généré sur le réseau, une première solution consiste à diminuer le nombre d’images par seconde. La meilleure solution consiste cependant à compresser les images, à l’instar de la voix.

Pour les données, aucun bit ne doit être altéré dans les phases de compression et de décompression : la décompression d’un fichier doit reproduire exactement le fichier original. Les taux de compression sont, en moyenne, de 1 pour 2, voire plus pour les images bitmap.

Mais, pour la vidéo, l'image décompressée peut être différente de la même image avant compression. Ainsi, une légère variation de couleur ou la suppression d'un pixel n'est pas perceptible lorsqu'ils affectent 1 image sur 25 par seconde. Ce principe permet d'atteindre des ratios de compression de 1 pour 100, voire 1 pour 300. Une perte de qualité est acceptable à partir du moment où elle est à peine perceptible.

La compression peut être **spatiale** (c'est-à-dire porter sur une image — mode intraframe) ou **temporelle** (c'est-à-dire porter sur plusieurs images dans le temps — mode interframe).

Dans le mode **interframe**, seule la différence entre deux images est transportée à partir d'une image de référence (*key frame*) ou à partir de l'image précédente. Une technique de compensation de mouvement est, de plus, utilisée : le delta est calculé non pas pixel par pixel mais par bloc de pixels en mouvement.

Dans le mode **intraframe**, l'image est divisée en blocs de 8×8 pixels : les pixels et leurs couleurs associées sont convertis en fréquence de changement de couleurs et amplitude des variations de couleurs. La moyenne est calculée sur le résultat, de manière que la perte d'information qui en résulte ne soit pas perceptible par l'œil. Les données sont ensuite compressées par un algorithme de type RLE (*Run Length Encoding*).

Par exemple, M-JPEG (*Motion JPEG*) utilise un algorithme spatial (série d'images JPEG), alors que MPEG utilise une combinaison des deux, et notamment l'algorithme DCT (*Discrete Cosine Transform*) pour le mode intraframe.

Standard	Bande passante requise	Remarque
MPEG1	1,5 Mbit/s	Compression préalable, puis décompression en temps réel
MPEG2	De 4 à 10 Mbit/s	Vidéo à la demande par des câblo-opérateurs pour HDTV. Codec en temps réel. Interfacé à AAL5.
MPEG4	64 Kbit/s à 2 Mbit/s	
M-JPEG		Agrégation d'images JPEG.
Cell B		Codec en temps réel. Algorithmes VQ (Vector Quantization) et RLE.
Cell A		Seule la décompression est effectuée en temps réel.
H.261	p x 64 Kbit/s (1 < p < 30) QCIF	Utilisé par les normes H.32x. Algorithme DCT.
H.263	p x 64 Kbit/s (1 < p < 30) Sub-QCIF et QCIF	Utilisé par les normes H.32x. Algorithme DCT.

Le CIF (*Common Intermediate Format*) est le format de base d'une image respectant les normes H.26x.

Format de l'image	Nombre de pixels pour la luminance	Nombre de lignes pour la luminance	Nombre de pixels pour la chrominance	Nombre de lignes pour la chrominance
sub-QCIF	128	96	64	48
QCIF	176	144	88	72
CIF	352	288	176	144
4CIF	704	576	352	288
16CIF	1 408	1 152	704	576

Généralement, le format QCIF est utilisé lorsque la bande passante est inférieure à 192 Kbit/s ($p \leq 3$).

La norme H.261 est semblable à MPEG, mais consomme moins de CPU pour la compression. De plus, l'algorithme permet d'augmenter la compression, au détriment cependant de la qualité des images qui bougent vite (s'il existe plus de 15 pixels de différence entre deux images, l'image devient floue).

Les problèmes posés par les transmissions audio et vidéo

Les flux audio et, dans une moindre mesure, les flux vidéo sont sensibles à plusieurs phénomènes :

- l'écho ;
- le délai de transit (encore appelé latence), c'est-à-dire le temps qui s'écoule entre la prononciation d'un mot et sa restitution côté récepteur ;
- la variation du délai de transit, appelée gigue (*jitter*, en anglais).

Le phénomène d'**écho** provient de la réflexion du signal sur le câble, surtout au niveau des convertisseurs 2 fils/4 fils présents en nombre sur le RTC. Il est surtout perceptible sur de longues distances et est amplifié si les délais de transit sont importants. Si ces derniers sont assez bas (moins de 50 ms), l'écho n'est pas perceptible et est masqué par la conversation. Étant donné que les délais de transit sont, la plupart du temps, supérieurs à 50 ms, des appareils spécifiques, appelés annulateurs d'écho (ITU G.165), sont systématiquement utilisés sur le RTC.

Le **délai de transit** affecte une conversation téléphonique : si le temps qui s'écoule entre la fin d'une phrase et sa réception complète par le récepteur est trop long, les personnes commencent à parler en même temps puis s'arrêtent en même temps, se coupent la parole, etc. La norme **G.114** préconise un délai de transmission maximal de 400 ms pour le RTC. Dans

les faits, la dégradation de la qualité de la voix est nettement perceptible lorsque les délais dépassent 300 ms. De plus, l'écho devient perceptible, ce qui ajoute à la mauvaise qualité.

Les **variations des délais de transit** (gigues) sont essentiellement dues à la charge du réseau qui doit traiter différents types de flux. Le RTC fonctionne en mode commutation de circuit : un canal de 64 Kbit/s est réservé pendant toute la durée de la conversation. Il y a donc très peu de gigue, voire pas du tout. L'inconvénient est que la bande passante du réseau n'est pas optimisée : les 64 Kbit/s restent accaparés même lorsqu'ils ne sont pas utilisés, pendant les silences ou en l'absence de transmission.

En revanche, un réseau de paquets, de trames ou de cellules véhicule simultanément une multitude de flux et exploite au mieux ses ressources en les partageant entre plusieurs utilisateurs. Le délai de traitement varie donc en fonction de sa charge.

La variation du délai crée ainsi des interruptions inattendues au milieu d'une phrase ou d'un mot (un paquet voix arrive plus tard que les autres), qui peuvent rendre une conversation inintelligible. Pour compenser la gigue, on utilise des tampons mémoire. Les trames arrivant en retard par rapport à la moyenne sont restituées immédiatement, mais celles arrivant en avance par rapport à cette même moyenne restent plus longtemps dans la mémoire tampon. L'inconvénient est donc que le délai de transit est augmenté proportionnellement à la taille du tampon. Généralement, sa taille correspond à un délai égal à environ deux fois celui du traitement du codec. Dans la pratique, les équipements font varier dynamiquement leur taille.

Concernant le téléphone, notons également une spécificité liée aux fonctionnalités offertes par le **DTMF** (*Dual-Tone MultiFrequency*). Ce signal permet à un utilisateur de dialoguer avec un serveur audiotel, par exemple pour accéder au rappel sur occupation (touche "5" avec France Télécom) ou pour consulter son compte bancaire. La compression de la voix empêche la transmission de ce type de signal. Celui-ci est donc codé, puis régénéré à son arrivée (G.729).

Estimation du temps de transit

Le délai de transit sur le réseau est la somme des délais induits par tous les équipements traversés : câbles, routeurs, commutateurs, etc.

Le délai de transit est donc la somme :

- du délai de sérialisation déterminé par la vitesse de la ligne et la taille du paquet ;
- des délais de traitement propres au codec ;
- du délai de transit dans un nœud (routeur ou passerelle), déterminé par l'empaquetage et le dépaquetage des données, augmentés des temps de traitement dus aux protocoles réseau (interprétation des en-têtes, routage, etc.).

La **sérialisation** est l'action d'envoyer les bits sur le câble. Plus le débit de la liaison est élevé, plus le temps de sérialisation est court. Par exemple, il faut 125 microsecondes pour envoyer un octet sur une LS à 64 Kbit/s contre 0,05 microsecondes sur une liaison à 155 Mbit/s. De même, plus les sites sont rapprochés, plus le délai de propagation est court. Par

exemple, il faut compter 32 ms de délai de propagation sur une LS 64 Kbit/s entre l'Europe et les États-Unis.

À titre d'exemple, le tableau suivant dresse la liste des délais nécessaires pour envoyer une trame de 30 ms générée par le codec G.723 (20 octets).

Opération	Délai
Traitement codec émission	7,5 ms (G.729)
Remplissage de la trame	30 ms de voix (20 octets)
Traitement codec réception	7,5 ms (G.729)
Tampon de gigue en réception	15 ms (deux fois le délai du codec)
Encapsulation IP/Frame-Relay	2 ms (traitement du routeur)
Sérialisation de 31 octets (20 + 5 + 6)* sur une ligne à 128 Kbit/s	3 ms
Sous-total	65 ms
Temps de transit réseau **	De 30 à 150 ms
Total	De 95 à 215 ms

(*) L'encapsulation dans un paquet IP ajoute 2 à 5 octets (avec compression des en-têtes RTP/UDP/IP, voir chapitre 15). Il faut ajouter à cela les en-têtes Frame-Relay (6 octets) ou ATM (5 octets) ou encore Ethernet, etc.

(**) Ce temps est égal à celui induit par les routeurs et/ou commutateurs et/ou passerelles, éventuellement augmenté de celui induit par la propagation des signaux sur les liaisons internes au réseau de l'opérateur.

Le transport des données multimédias

À la différence des données applicatives, la voix et la vidéo — ainsi que, dans une moindre mesure, la télécopie — acceptent l'altération des données. Il en résulte une dégradation de la qualité du son, de l'image ou de la page, qui peut toutefois être tolérée par les personnes qui les reçoivent. En plus du délai de transit et de sa variation, des paquets peuvent donc être perdus, mais bien sûr dans une certaine limite.

Qualité	Voix	Fax	Vidéo
Bonne (dégradations non perceptibles)	Délai < 200 ms Gigue < 15 ms Perte < 5 %	< 200 ms Gigue < 100 ms Perte < 2 %	< 200 ms Gigue < 15 ms Perte < 5 %
Limite de l'acceptable	Délai = 400 ms Gigue = 30 ms Perte = 10 %	Délai = 300 ms Gigue < 1 000 ms Perte = 4 %	Délai = 400 ms Gigue = 20 ms Perte = 10 %

En outre, les liens réseau doivent être correctement dimensionnés, afin de supporter les débits générés par les flux multimédias. Ces derniers sont variables, car ils dépendent de la qualité voulue pour le son et l'image.

Application	Bande passante requise
Audio (voix qualité téléphone)	64 Kbit/s sans compression (<i>toll quality</i>) De 11 à 30 Kbit/s selon la compression
Vidéoconférence	384 Kbit/s (qualité visio professionnelle) 10 Mbit/s (qualité DVD)
Tableau partagé	De 10 à 64 Kbit/s

Sur le réseau téléphonique à commutation de circuit (RTC, RNIS), la voix occupe un débit fixe réservé entre l'appelant et l'appelé pendant toute la durée de la communication. Sur un réseau à commutation de paquets, tel qu'IP, aucun débit n'est réservé, et celui-ci peut varier dans le temps (par exemple, les silences ne sont pas transmis, etc.).

La voix peut ainsi être transmise sur IP — on parle alors des produits VoIP (*Voice Over Internet Protocol*) —, directement sur Frame Relay (VoFR) ou sur ATM (VoATM).

Dans tous les cas, le transport de la voix et de la vidéo requiert des mécanismes spécifiques, afin d'assurer la **qualité de service** requise (débit, délai de transit, variation du délai de transit, perte de paquets tolérée).

Un autre problème, qui cette fois n'est pas nouveau, est de dimensionner les liens réseau en fonction du nombre d'utilisateurs. On parlera ici de canaux, un canal correspondant à une session voix ou vidéo.

Une règle simple, mais peu économique, est de prévoir autant de canaux que d'utilisateurs : le débit total est alors égal au débit d'un canal VoIP multiplié par le nombre d'utilisateurs. Une règle plus complexe, mais plus rationnelle, consiste à s'appuyer sur le nombre de communications simultanées et à utiliser des modèles statistiques permettant de prévoir le nombre de canaux nécessaires en fonction de différents paramètres, tels que le nombre d'utilisateurs, la durée moyenne des communications, le taux de débordement acceptable, etc.

Les modèles standards s'appuient sur des distributions de Poisson qui permettent de convertir la nature aléatoire des appels en probabilité, sur une unité de mesure appelée Erlang ou CCS (*Centum Call Seconds*).

En conclusion, un travail d'ingénierie important est nécessaire pour préparer son réseau IP au multimédia.

13

Le routage des flux multimédias

Qui dit multimédia dit également diffusion d'un flux audio ou vidéo à plusieurs destinataires, dans le cadre d'une conférence, par exemple. Un même film vidéo peut ainsi être dupliqué en autant de flux qu'il y a de destinataires.

Afin de limiter la charge induite sur le réseau, il est bien plus judicieux de dupliquer ce flux au plus près des destinataires. Pour ce faire, trois mécanismes doivent être mis en place sur notre réseau IP :

- un adressage permettant de désigner des groupes de machines plutôt qu'une seule ;
- un mécanisme permettant d'identifier les groupes actifs au sein du réseau ;
- un mécanisme permettant de router les paquets en les dupliquant le moins possible.

Sans cela, notre réseau IP serait vite engorgé, surtout sur les liaisons WAN pour lesquelles le débit est compté.

Dans ce chapitre, vous apprendrez ainsi :

- à utiliser l'adressage multicast ;
- à gérer les groupes de diffusion ;
- à configurer les algorithmes de routage multicast DVMRP, MOSFP et PIM ;
- à choisir l'un de ces algorithmes en fonction de vos besoins.

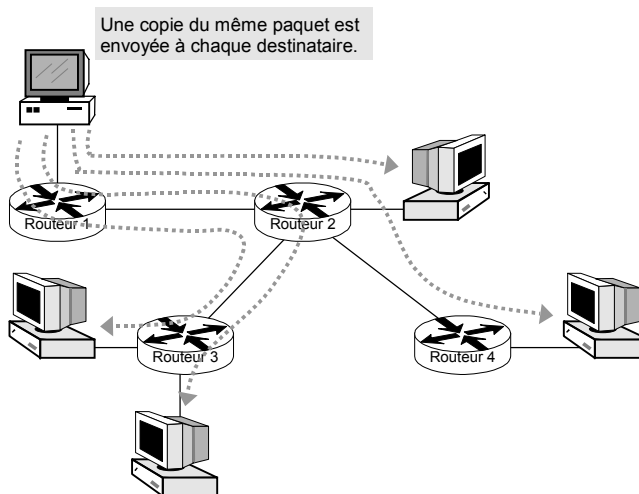
La diffusion sur un réseau IP

Un participant à une audioconférence parle à tous les autres participants : sa voix est numérisée, puis découpée en paquets IP qui sont ensuite envoyés sur le réseau.

Au premier abord, il est possible de transmettre ces paquets sur un réseau IP classique : une copie de ce paquet est alors transmise à chaque destinataire. Chacun d'eux est, en effet, identifié par une adresse IP unique qui est insérée dans le champ destination du paquet. Les routeurs se servent de cette adresse pour acheminer le paquet jusqu'au destinataire.

Figure 13-1.

Diffusion des paquets unicast.



Cet exemple montre que ce type de fonctionnement n'est pas adapté à une diffusion : le réseau est inondé par des paquets dupliqués dès la source d'émission. Les adresses utilisées (classes A, B ou C) sont, en effet, de type **unicast**, car à une adresse est associé une machine cible. Par extension, les paquets sont dits unicast.

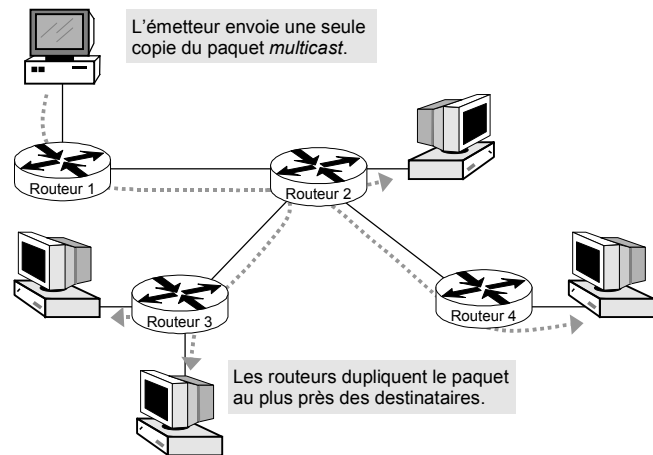
L'adressage IP propose une autre classe d'adresses, la classe D. Ces adresses sont dites **multicast**, car elles désignent un groupe de machines. Il ne s'agit pas d'une adresse de broadcast, car seules les machines qui sont configurées pour accepter une adresse multicast prendront en compte le paquet.

La plage réservée pour la classe D s'étend de 224.0.0.0 à 239.255.255.255. L'adresse 224.0.0.0 n'est attribuée à aucun groupe ; l'adresse 224.0.0.1 permet d'adresser toutes les machines sur un réseau (le réseau local sur lequel est émis le paquet). L'adresse 224.0.0.2 permet d'adresser, plus spécifiquement, tous les routeurs sur un réseau.

Des adresses de groupes permanents sont attribuées officiellement par l'IANA (*Internet Assigned Number Authority* — www.iana.org — RFC 1700).

Groupe concerné	Exemple d'adresse attribuée
Toutes les machines sur le réseau local	224.0.0.1
Tous les routeurs sur le réseau local	224.0.0.2
Tous les routeurs DVMRP	224.0.0.4
Tous les routeurs OSPF	224.0.0.5
Tous les routeurs OSPF désignés	224.0.0.6
Tous les routeurs RIP	224.0.0.9
Tous les routeurs PIM	224.0.0.13
Messages RSVP encapsulés dans UDP	224.0.0.14
NTP (<i>Network Time Protocol</i>)	224.0.1.1
Artificial Horizons — Aviator	224.0.1.5
Music-Service	224.0.1.16
IETF-2-Video	224.0.1.15
Microsoft-ds	224.0.1.24

Figure 13-2.
Diffusion d'un paquet multicast.



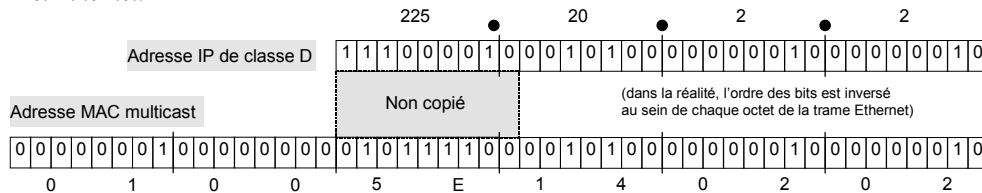
Sans les mécanismes du multicast, le serveur devrait, tout d'abord, déterminer les destinataires (leur adresse IP ou leur nom), puis envoyer autant d'exemplaires du paquet qu'il y a de destinataires.

L'adresse IP source est toujours celle de la station qui envoie le paquet (une adresse unicast). Aucun paquet n'est émis avec une adresse source multicast. Une adresse de classe D identifie toujours un groupe de destinataires.

Sur les réseaux qui prennent en charge ce type d'adressage, le multicast IP utilise les fonctions de multicast du niveau 2, c'est-à-dire des trames multicast. Sur un réseau Ethernet, les paquets multicast IP sont envoyés dans des trames multicast dont l'adresse MAC de destination commence par "01-00-5E". Les 23 derniers bits de cette adresse correspondent aux 23 derniers bits de l'adresse IP.

Figure 13-3.

Correspondance entre les adresses multicast IP et Ethernet.



Un PC envoie donc tous ses paquets multicast IP dans une trame multicast Ethernet. Le mécanisme de routage par défaut (voir chapitre 11) ne s'applique, en effet, qu'aux paquets unicast.

Sur les réseaux ne disposant pas de la fonction multicast de niveau 2 (X.25, Frame Relay, ATM, par exemple), les paquets multicast sont transportés dans les trames unicast de niveau 2.

La gestion des groupes de diffusion

Un groupe de diffusion (groupe *multicast*) est dynamique : ses membres peuvent adhérer au groupe ou le quitter à tout moment, être dispersés à travers le monde, adhérer à plusieurs groupes en même temps. Aucune restriction quant au nombre de participants n'est également appliquée. Le groupe peut être permanent ou non. Un participant peut être actif ou non (la machine est éteinte).

La première tâche pour un participant (une machine connectée sur le réseau) est donc de se faire connaître. Pour cela, il dispose du protocole **IGMP** (*Internet Group Membership Protocol*) défini par la RFC 1112, datée de 1989, et mis à jour en 1997 par la RFC 2236 (IGMP v2).

En principe, toutes les piles TCP/IP, et en particulier celle de Windows, supportent IGMP. Aucune configuration spécifique n'est nécessaire, car la gestion des groupes n'est pas réalisée manuellement mais directement par les applications *via* des API (*Application Programming Interface*).

Par exemple, l'interface Winsock offre les primitives permettant d'entrer et de sortir d'un groupe :

- *JoinHostGroup* (adresse IP du groupe, numéro de l'interface réseau) ;
- *LeaveHostGroup* (adresse IP du groupe, numéro de l'interface réseau).

Ces fonctions permettent à la pile TCP/IP d'accepter d'une part les paquets dont l'adresse IP de destination lui correspond et, d'autre part, les paquets dont l'adresse IP de destination est celle du ou des groupes multicast auxquels l'interface réseau est non jointe. La pile TCP/IP maintient ainsi une liste d'appartenance pour chaque interface réseau de la machine.

Ainsi, plusieurs logiciels utilisent déjà le multicast à votre insu :

- les serveurs WINS de Windows NT, qui se placent d'office dans le groupe 224.0.1.24 qui est utilisé pour dialoguer avec des partenaires de réplication ;
- le logiciel de visioconférence Netmeeting ;
- le logiciel de diffusion audio et vidéo RealG2 Player.

La plupart des routeurs prennent également en charge ce protocole. La configuration d'un routeur Cisco consiste simplement à activer la fonction multicast d'IP :

```
int e 0
ip multicast-routing
```

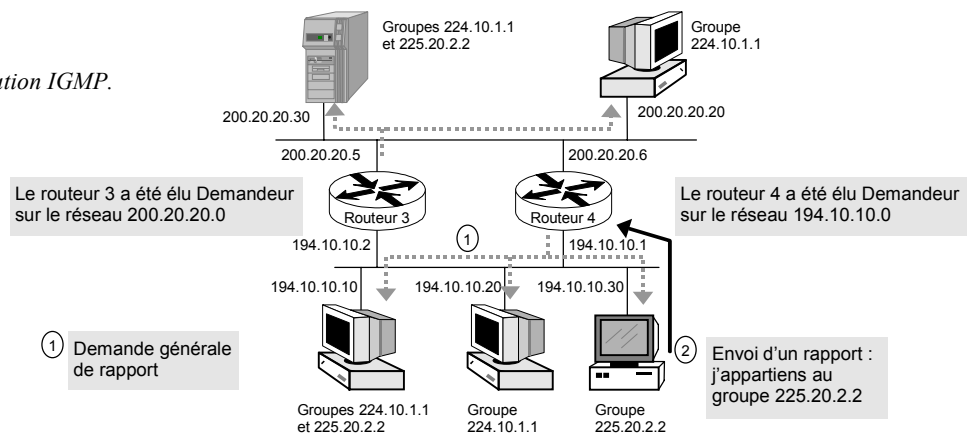
Si le routeur dispose de plusieurs interfaces sur le même réseau local, une seule doit être configurée avec IGMP.

On peut indiquer au routeur de devenir membre d'un groupe. Cela permet aux exploitants de savoir si un groupe est joignable en utilisant la commande *ping*. Si personne n'est actif sur ce groupe, le routeur pourra au moins répondre au *ping* :

```
ip igmp join-group 224.10.1.1
```

Considérons l'exemple suivant : chacun des deux routeurs est élu Demandeur pour un réseau en fonction de son adresse IP.

Figure 13-4.
*Exemple
de configuration IGMP.*



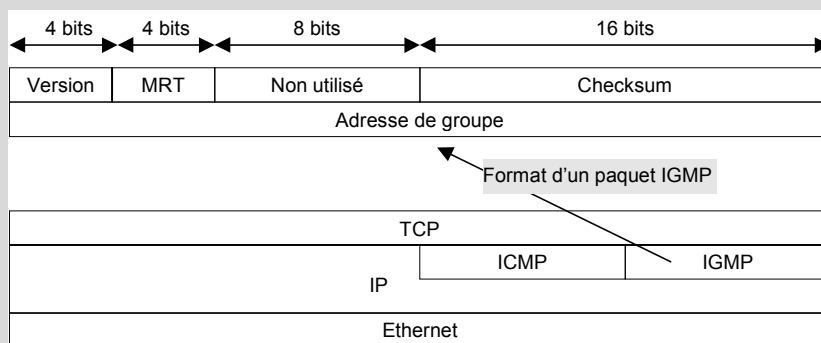
LE POINT SUR IGMP (RFC 1112 ET 2236)

IGMP (*Internet Group Membership Protocol*) est un protocole qui fonctionne conjointement avec IP, au même titre qu'ICMP. Il permet aux membres d'un groupe de signaler leur présence au routeur le plus proche, celui connecté au réseau local. Si plusieurs routeurs sont connectés à un même réseau local, celui qui a la plus petite adresse IP est élu Demandeur.

Le routeur Demandeur émet périodiquement une **demande générale de rapport** (adresse de destination 224.0.0.1). Les membres de tous les groupes actifs sur ce réseau local lui répondent en envoyant un **rapport** pour chaque groupe auquel ils appartiennent.

Le routeur maintient à jour une liste des groupes dont au moins un membre est actif. Les routeurs non demandeurs n'émettent pas de demande de rapport, mais lisent les rapports et mettent à jour leur table.

Il existe trois types de paquets IGMP : demande de rapport, sortie d'un groupe (demandée par un membre) et rapport d'activité. Un quatrième type, le rapport IGMP version 1, est supporté à des fins de compatibilité.



Le champ MRT (*Max Response Time*) est utilisé dans les paquets de demande de rapport pour indiquer, en dixièmes de secondes, le délai maximal autorisé pour envoyer un rapport d'activité. Au-delà de ce délai, le destinataire est considéré comme n'étant pas actif dans le groupe.

Un membre peut également demander à sortir du groupe (message à destination de 224.0.0.2). Le routeur Demandeur émet alors une **demande spécifique de rapport** (adresse de destination identique à celle du groupe) pour s'assurer qu'un membre au moins est encore actif.

Les routeurs sont à l'écoute de tout message multicast, et prennent donc en compte tous les rapports. La commande suivante montre le résultat obtenu sur un routeur Cisco :

```
Routeur3# show ip igmp interface
Ethernet0 is up, line protocol is up
```

```

Internet address is 194.10.10.2, subnet mask is 255.255.255.0
IGMP is enabled on interface
IGMP query interval is 120 seconds
Inbound IGMP access group is not set
Multicast routing is enabled on interface
Multicast TTL threshold is 0
Multicast designated router (DR) is 194.10.10.1
Multicast groups joined: 224.10.1.1 225.20.2.2
Ethernet1 is up, line protocol is up
Internet address is 200.20.20.5, subnet mask is 255.255.255.0
IGMP is enabled on interface
IGMP query interval is 120 seconds
Inbound IGMP access group is not set
Multicast routing is enabled on interface
Multicast TTL threshold is 0
Multicast designated router (DR) is 200.20.20.5
Multicast groups joined: 224.10.1.1 225.20.2.2

```

Le routeur 3 n'est pas Demandeur pour le réseau 194.10.10.0 ; il n'envoie donc pas de requête. En revanche, il prend en compte tous les rapports qu'il voit passer :

```

Routeur4# show ip igmp groups
IGMP Connected Group Membership

```

Group Address	Interface	Uptime	Expires	Last Reporter
224.10.1.1	Ethernet0	17:15:15	0:02:05	194.10.10.20
224.10.1.2	Ethernet1	1:01:01	0:01:05	220.20.20.30
224.10.1.1	Ethernet1	17:15:40	0:01:40	220.20.10.30
225.20.2.2	Ethernet0	1:00:45	0:01:50	194.10.10.10

De même, le routeur 4 n'est pas Demandeur pour le réseau 220.20.20.0 ; il n'envoie donc pas de requête. En revanche, il prend également en compte tous les rapports qu'il voit passer.

La périodicité d'envoi des demandes générales de rapports peut être paramétrée comme suit :

```
ip igmp query-interval 120
```

← Toutes les 120 secondes

La version 3 d'IGMP, en cours d'étude, permettra aux machines d'indiquer au routeur l'adresse IP source pour laquelle elle accepte de recevoir des paquets multicast ; ainsi, la diffusion du paquet prendra en compte l'adresse de l'émetteur en plus du groupe destination.

Le routage des flux multicast

Le protocole IGMP permet aux routeurs de détecter les groupes situés sur leurs réseaux locaux (les réseaux auxquels une interface est connectée).

Mais les routeurs ne savent pas où sont situés les autres membres du groupe, puisque IGMP n'a qu'une portée locale. Pour cela, il faut utiliser des protocoles de routage spécifiques au multicast. Trois standards sont disponibles :

- DVMRP (*Distance Vector Multicast Routing Protocol*), analogue au protocole RIP adapté au multicast ;
- MOSPF (*Multicast Open Shortest Path First*), une extension d'OSPF ;
- PIM (*Protocol Independent Multicast*), spécialement dédié au multicast.

Le choix de l'un de ces protocoles dépend de nombreux paramètres. Les paragraphes suivants décrivent donc leurs principes de fonctionnement, afin de mieux cerner leurs conséquences sur l'architecture de notre réseau.

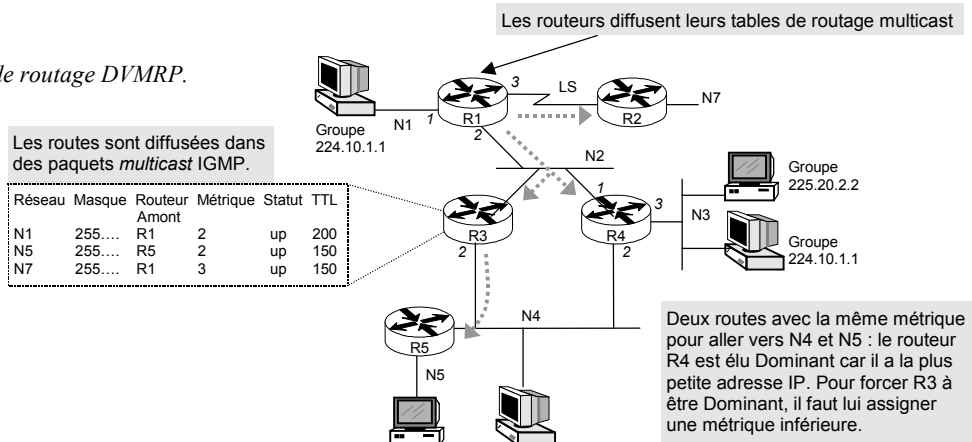
Le routage à l'aide de DVMRP

Le premier protocole de routage multicast a été DVMRP (*Distance Vector Multicast Routing Protocol*). Comme son nom l'indique, il repose sur un algorithme de calcul du plus court chemin, basé sur le plus petit nombre de routeurs à traverser pour atteindre une destination (nombre de sauts, appelé métrique).

Le principe est en cela identique à RIP (*Routing Information Protocol*) : les routeurs s'échangent l'intégralité de leurs tables de routage. DVMRP doit donc être utilisé en supplément d'un protocole de routage unicast (RIP, OSPF, etc.).

Figure 13-5.

Diffusion des tables de routage DVMRP.



La commande suivante permet d'activer DVMRP sur notre routeur R1, qui est un Netbuilder de marque 3com. Elle doit être utilisée pour chaque interface gérant le trafic *multicast*. La commande MIP permet d'activer le protocole IGMP :

```
setdefault -MIP control = enable
setdefault !1 -DVMRP control = enable
setdefault !2 -DVMRP control = enable
setdefault !3 -DVMRP control = enable
```

Lorsque plusieurs routeurs coexistent sur le même réseau local, seul l'un d'eux a la charge de diffuser les paquets multicast, afin d'éviter la duplication des paquets. Comme pour RIP, la route choisie repose sur le plus petit nombre de sauts (la métrique). En cas d'égalité de métrique, le routeur dominant élu est celui qui possède la plus petite adresse IP.

Pour forcer R3 à être le routeur dominant sur le réseau N4, il faut lui attribuer une métrique inférieure à celle de R4 :

```
#Routeur R3
setdefault !2 -DVMRP metric = 5
#Routeur R4
setdefault !2 -DVMRP metric = 10
```

La commande suivante permet de visualiser la table de routage du routeur R4 :

```
show -dvmrp routetable long
```

SourceSubnet	SubnetMask	FromGateway	Metric	Status	TTL	InPort	OutPorts
200.10.10.0	255.255.255.0	200.20.20.1	2	Up	200	1	2 3*
200.30.30.0	255.255.255.0		1	Up	150	3	1 2
...							

La colonne "FromGateway" indique le routeur le plus proche qui mène à la source (un champ vide indique que le réseau est directement connecté au routeur R4).

La colonne "InPort" (*Incoming Port*) indique l'interface par laquelle arrivent les paquets multicast émis par la source précisée dans la colonne "SourceSubnet".

La colonne "OutPorts" (*Outgoing Ports*) donne la liste des ports vers lesquels seront diffusés par défaut les paquets multicast issus de la source indiquée dans la colonne "SourceSubnet". Un astérisque indique que le port conduit à une feuille de l'arbre, c'est-à-dire qu'aucun routeur ne se trouve en dessous.

En plus de la table de routage, le routeur gère une table de diffusion construite lorsque les premiers paquets *multicast* transitent par le routeur. Elle permet d'enregistrer les groupes identifiés pour chaque réseau source.

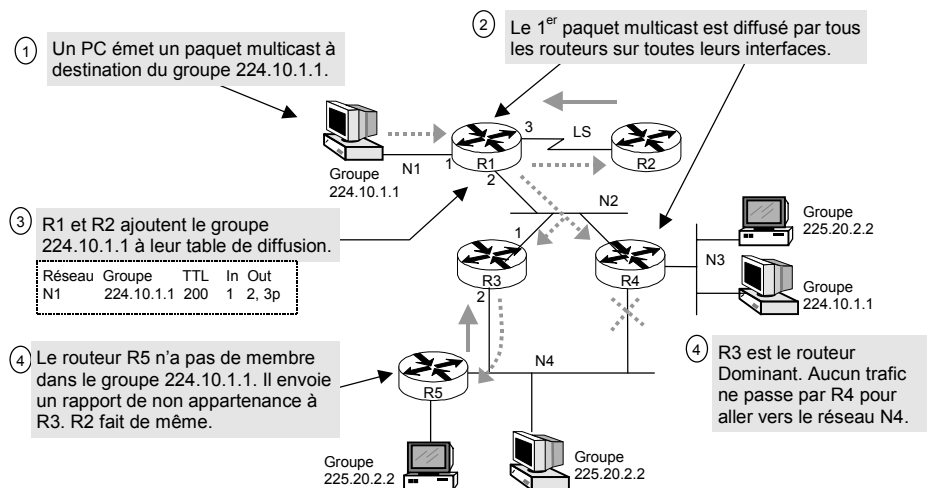
La commande suivante affiche la table de diffusion pour le routeur R4 :

```
show -dvmrp forwardtable
```

```
SourceSubnet MulticastGroup TTL InPort OutPorts
200.10.10.0 224.10.1.1 200 1 1 2p
200.30.30.0 224.10.1.1 150 3 2p 3
          225.20.2.2 200 3 2p 3
...
```

La colonne “ MulticastGroup ” indique la liste des groupes de diffusion des paquets en provenance de la source indiquée dans la colonne “ SourceSubnet ”. Les autres colonnes ont la même signification que celles de la table de routage. L’indicateur “ p ” (*prune*) indique qu’un message de non-appartenance a été reçu ; par conséquent, aucun paquet multicast ne sera envoyé vers cette interface. Dans notre cas, cela indique que R3 est dominant.

Figure 13-6.
Routage
d'un paquet
multicast.



Dans notre réseau, le routeur R1 a diffusé le paquet multicast vers les interfaces 2 et 3. Il met alors à jour sa table de diffusion. L’interface 3 a par la suite été marquée *p* (*pruned*) car le routeur R2 lui a renvoyé un rapport de non appartenance. En effet, ce dernier n’a détecté (via IGMP) aucun membre actif pour le groupe 224.10.1.1.

De même, le routeur R5 renvoie un message de non-appartenance au routeur R3, qui ne lui transmettra alors plus aucun paquet pour le groupe 224.10.1.1. Le routeur R3 fait ensuite de même vis-à-vis de R1 et R4.

LE POINT SUR DVMRP (RFC 1075)

DVMRP (*Distance Vector Multicast Routing Protocol*) utilise son propre protocole de routage de paquets multicast, qui est analogue à celui de RIP : les routeurs s'échangent l'intégralité de leurs tables de routage et calculent les routes sur la base d'une **métrique** (nombre de sauts en termes de routeurs).

Il existe plusieurs implémentations de DVMRP qui diffèrent par l'algorithme utilisé pour construire la table de diffusion : **TRPB** (*Truncated Reverse Path Broadcasting*) ou **RPM** (*Reverse Path Multicasting*), ce dernier étant le plus répandu car plus performant. Les routeurs 3com et le démon Unix *mrouted* (à partir de la version 3.8) utilisent RPM.

Le principe du TRPB est le suivant : un routeur recevant un paquet multicast le transmet sur toutes les autres interfaces. Les routeurs situés en aval reçoivent donc le paquet. S'ils ne disposent d'aucun membre déclaré ni d'aucun autre routeur en aval, ils renvoient un rapport de non-appartenance au routeur situé en amont. Ce dernier ne leur transmettra alors plus de paquets multicast.

Si un nouveau membre s'enregistre sur un des routeurs situés en aval, celui-ci enverra au routeur situé en amont un message d'annulation, pour recevoir à nouveau les paquets *multicast* destinés au groupe en question.

Le principe du RPM reprend celui du TRPB, mais pousse plus loin la remontée d'information : un routeur qui ne dispose pas de membre ni de routeur en aval ayant de membre envoie un rapport de non-appartenance aux routeurs situés en amont, de sorte que lui-même ne reçoive pas de paquets multicast. Le rapport peut ainsi remonter jusqu'à la source si nécessaire.

Un paquet DVMRP est composé d'un en-tête IGMP auquel sont jointes des données de longueur variable (512 octets au maximum) formatées sur le mode : " Commande, Valeur " ou " Commande, nombre de valeurs, valeur 1, valeur 2, etc. ". L'exemple suivant montre un paquet d'annonce de la route 194.10.10.0.

4 bits		4 bits	8 bits	16 bits	
Version	Type	Sous-type		Checksum	
Cmde = AFI		Valeur = IP		Cmde = Métrique	Valeur = 2
Cmde = Intinity		Valeur = 6		Cmde = Subnet	Compte = 1
Valeur = 255.255.255.0					
Cmde = DA		Compte = 1		Valeur = 194.10	
.10.0					

Le champ Cmde contient le code d'une commande qui détermine la taille et la signification du champ Valeur. Plusieurs commandes se suivent dans un paquet. Le champ Sous-type détermine le type de message : requête, réponse, rapport de non-appartenance ou annulation d'un rapport de non-appartenance.

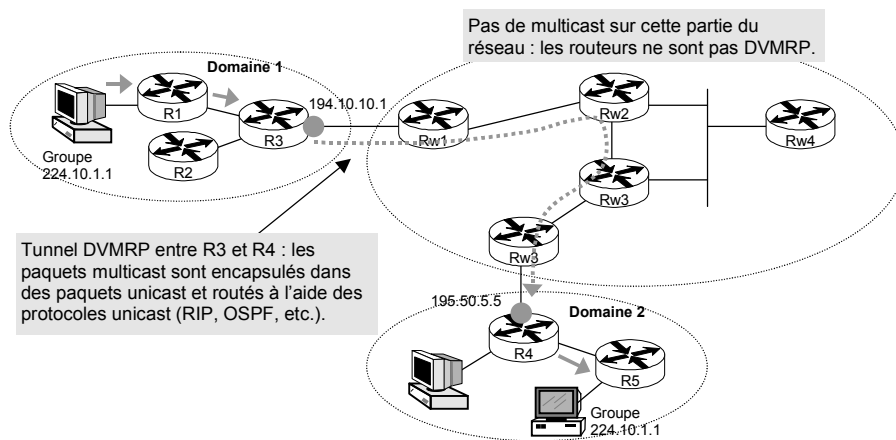
La commande AFI est toujours présente : elle indique simplement que la famille d'adresses est IP (seul supporté). La commande DA indique une adresse destination.

L'inconvénient de DVMRP est que les paquets multicast doivent périodiquement être envoyés aux routeurs situés en aval pour tenir compte d'éventuels changements de topologie ou d'appartenance à un groupe (pour un groupe donné, les messages d'annulation ne remontent que si un multicast a été précédemment reçu et un message de non-appartenance émis). Il en résulte une nouvelle cascade de rapports de non-appartenance ou d'annulation remontant vers les routeurs situés en amont.

Sur un réseau, tous les routeurs ne sont pas DVMRP. La mise en place du routage multicast est, en effet, progressive. En outre, le fait de ne pas installer de routeurs DVMRP partout permet de définir des domaines dans le but de circonscrire la diffusion des paquets multicast.

Pour assurer néanmoins la diffusion des paquets multicast, il est possible de créer un tunnel entre deux routeurs DVMRP séparés par des routeurs qui ne savent pas router les paquets multicast.

Figure 13-7.
*Principe
du tunneling
DVMRP.*



Sur nos routeurs 3com, la création du tunnel IP est réalisée simplement en indiquant les adresses des deux routeurs :

#Sur le routeur R3

```
setdefault !1 -DVMRP mon_tunnel = 194.10.10.1 195.50.5.5
```

#Sur le routeur R4

```
setdefault !2 -DVMRP mon_tunnel = 195.50.5.5 194.10.10.1
```

#Activation de DVMRP sur le tunnel

```
setdefault mon_tunnel -DVMRP control = enable
```

Afin d'atténuer les effets d'un trafic multicast important sur un réseau non dimensionné pour cela, il est possible d'en limiter le débit à 64 Kbit/s, par exemple :

```
# Sur les routeurs R3 et R4
setdefault mon_tunnel -dvmrp ratelimit = 64
```

De même, il est possible d'augmenter la périodicité d'échange des tables de routage entre les deux routeurs (60 secondes par défaut) :

```
setdefault -dvmrp updatetime = 120
```

Enfin, la durée de validité des entrées dans la table de diffusion peut également être allongée (dans notre cas, 3 600 secondes) :

```
setdefault -dvmrp cachetime = 3600
```

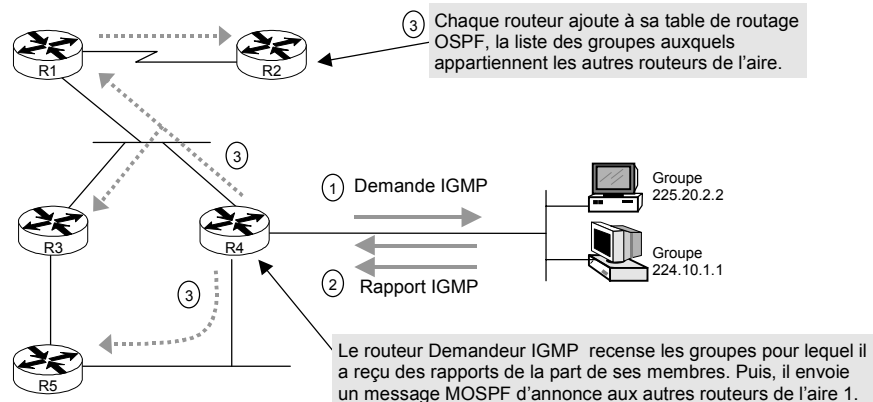
Le routage à l'aide de MOSPF

Le protocole MOSPF (*Multicast Open Shortest Path First*) est une extension d'OSPF permettant de router les paquets IP *multicast*. La détection des membres actifs d'un groupe est, comme toujours, assurée par IGMP, mais MOSPF permet aux routeurs de diffuser les groupes auxquels ils appartiennent (c'est-à-dire pour lesquels au moins un membre est actif).

Reprenons l'exemple de notre réseau : chaque routeur est configuré avec IGMP et MOSPF. Localement, chaque routeur établit une liste de tous les groupes pour lesquels il existe un membre actif sur son (ou ses) interface LAN.

En même temps, les routeurs s'échangent des messages d'annonce d'état des liens, suivant en cela le fonctionnement classique d'OSPF (voir chapitre 11). Cela leur permet de connaître la topologie du réseau.

Figure 13-8.
*Diffusion
d'un paquet
multicast
par MOSPF.*



Dès le premier enregistrement d'un membre de groupe *via* IGMP, les routeurs sont considérés par OSPF comme étant membres du groupe. Ces routeurs s'échangent alors des messages d'annonce d'appartenance à un groupe. Tous les routeurs savent désormais quel routeur appartient à quel groupe.

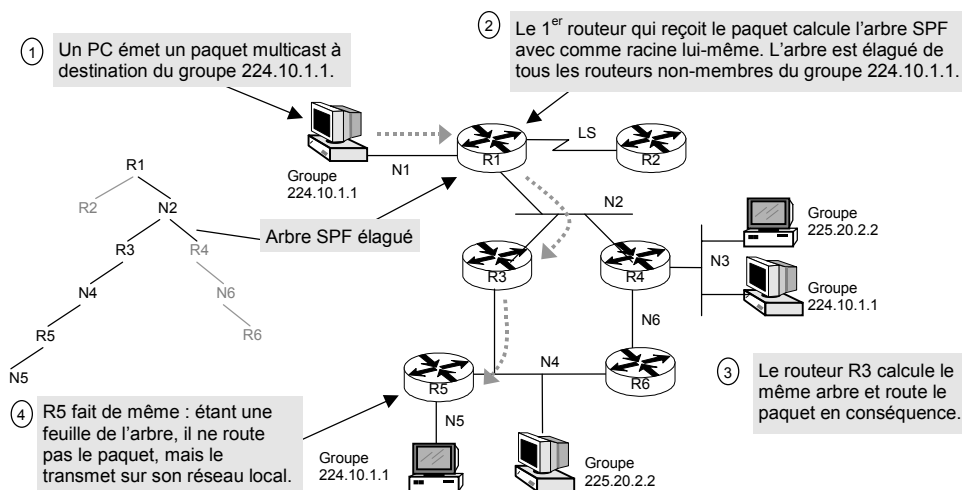
Sur nos routeurs 3com, l'activation de MOSPF est effectuée grâce à la commande suivante :

```
#Activation d'igmp
setdefault -mip control = enable

#Activation de mospf sur chaque interface
setdefault !1 -mospf control = enable
setdefault !2 -mospf control = enable
```

Le routage des paquets *multicast* au sein d'une aire est réalisé en fonction de la source (adresse IP *unicast* de la machine), de la destination (adresse IP *multicast* du groupe) et du TOS (*Type of Service*). Le routeur construit pour cela un arbre SPF (*Shortest Path First*) dont tous les routeurs non-membres du groupe indiqué dans le paquet ont été supprimés. Cet arbre est calculé à la demande lorsqu'un paquet *multicast* arrive sur le routeur.

Figure 13-9.
Calcul de l'arbre
SPF élagué.



Dans cet exemple, il est à noter que le routeur Demandeur IGMP doit être un routeur MOSPF, seul capable d'envoyer des messages d'annonce d'appartenance à un groupe. Il faut donc affecter une priorité supérieure aux routeurs MOSPF pour qu'ils soient élus routeurs désignés (au sens OSPF).

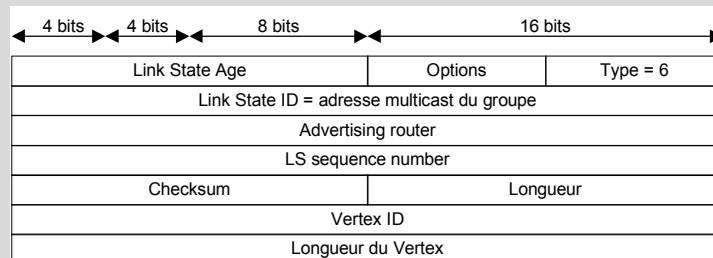
LE POINT SUR MOSPF (RFC 1584)

Au sein d'une aire (voir encadré "Le point sur OSPF"), les routeurs s'échangent des **messages d'annonce d'appartenance à un groupe**, permettant à chacun d'eux de disposer d'une visibilité complète de la topologie du réseau.

Chaque routeur calcule un **arbre SPF** pour chaque triplet adresse source unicast/adresse destination *multicast*/TOS (*Type of Service*). Ce calcul est effectué à la demande, c'est-à-dire lorsque le premier paquet multicast arrive. L'arbre ainsi calculé est conservé dans une mémoire cache, puis détruit au bout d'un certain temps lorsqu'aucun autre paquet de ce type n'a été reçu.

L'arbre est **élagué** : tous les routeurs qui ne sont pas membres du groupe indiqué dans le paquet sont supprimés. Les réseaux terminaux (*stub*) n'ont pas besoin d'être pris en compte : la diffusion locale au routeur est, en effet, assurée par IGMP.

Un chemin *multicast* est donc calculé en construisant un arbre élagué du plus court chemin dont la racine est l'émetteur du paquet (la RFC emploie l'expression : "*pruned shortest-path tree rooted at the packet's IP source*"). Pour la diffusion des appartenances à un groupe, un nouveau paquet d'annonce a été ajouté : *group-membership-LSA*.



Par ailleurs, les modifications suivantes ont été apportées à OSPF :

- Le champ "Option" des paquets Hello, Description de la base et Annonce d'état de lien contient un nouvel indicateur, le bit MC, qui indique si le routeur prend en charge l'extension multicast d'OSPF (MOSPF).
- Le champ "Rtype" des paquets d'annonce d'état de liens contient un nouvel indicateur, le bit W, qui indique si le routeur accepte ou non les multicast provenant de n'importe quelle source.

Un routeur calcule autant d'arbres qu'il y a de participants à une conférence, et calcule autant de variantes de cet arbre qu'il y a de groupes destinataires. La prise en compte du TOS augmente encore le nombre de combinaisons.

Le nombre d'arbres peut donc être très important. C'est pour cela qu'ils sont calculés à la demande et que le résultat est conservé en mémoire cache. Cela implique également que, plus le nombre de machines et de groupes est élevé, plus la CPU des routeurs est sollicitée pour calculer les arbres.

Le résultat du calcul d'un arbre est conservé aussi longtemps qu'un trafic entre le couple d'adresses IP *unicast* source/*multicast* destination existe, et jusqu'à ce qu'un changement de topologie intervienne. La taille mémoire requise est donc plus importante par rapport à un routeur OSPF classique (14 à 20 octets par triplet adresses source/groupe/TOS).

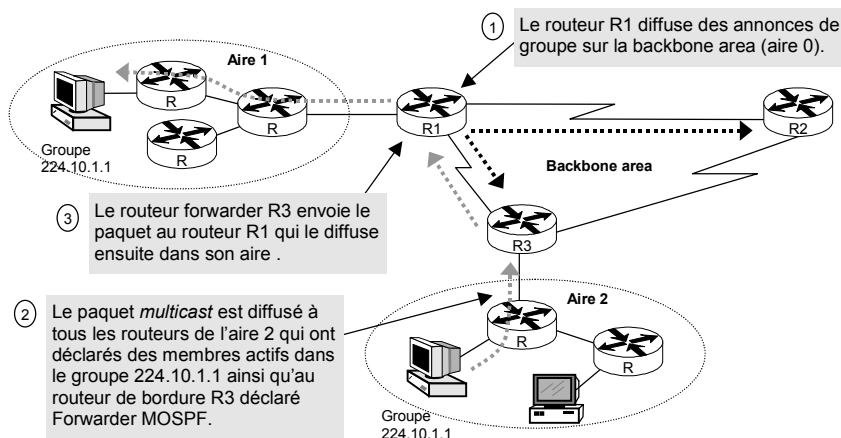
Le dimensionnement mémoire/CPU dépend du constructeur du routeur (architecture et puissance), de l'importance du réseau (nombre de machines *multicast* et nombre de groupes) ainsi que des performances visées (utilisation intensive ou non des téléconférences, utilisation du réseau à d'autres fins que le multimédia, etc.).

Il est également recommandé de limiter le nombre de TOS, voire de ne pas activer cette fonctionnalité, afin de limiter le nombre de combinaisons d'arbres SPF.

La diffusion des paquets multicast entre aires et systèmes autonomes utilise le principe appelé Multicast Forwarder : les routeurs de bordure et les routeurs inter-AS déclarés comme tels ne sont jamais supprimés des arbres SPF. En conséquence, tous les paquets multicast leur seront envoyés. De leur côté, ces routeurs calculent un arbre SPF pour chaque aire (ou AS) à laquelle ils sont rattachés.

Figure 13-10.

*Routing
des paquets multicast
entre aires.*



Les annonces d'appartenance à un groupe sont diffusées au sein de l'aire 0 (*backbone area*) par tous les routeurs de bordure. Tous les routeurs participant à la backbone area connaissent donc tous les groupes présents au sein du système autonome. Les informations provenant de la backbone area ne sont cependant pas transmises aux autres aires. Seuls les Forwarders conservent ces informations.

L'architecture et le dimensionnement des réseaux constituant la *backbone area* sont donc très importants. Il faut tenir compte du nombre de postes de travail et de leur répartition entre les aires. L'aire backbone concentre, en effet, la somme des flux circulant dans chacune des aires où les groupes sont actifs. Si cela est possible, il est donc conseillé de limiter le nombre d'aires ou bien de répartir des groupes entre plusieurs aires.

La configuration pour déclarer Forwarder un routeur de bordure 3com est la suivante :

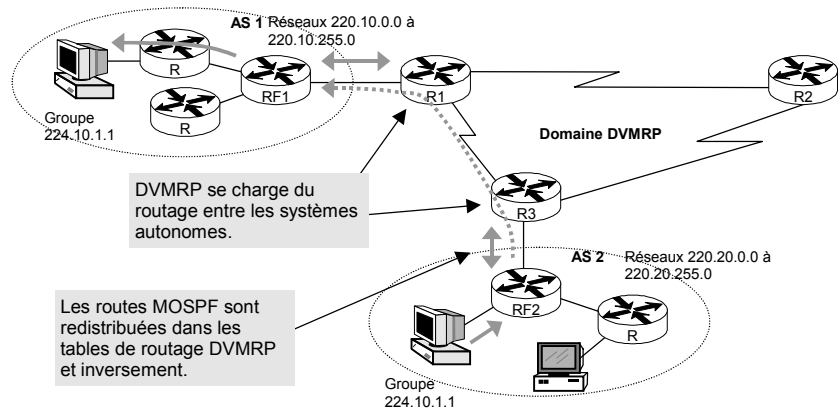
```
SETDefault -MOSPF MABR = Enable
```

Le mot clé MABR signifie *Multicast Area Border Router*.

Le principe est identique pour le routage entre systèmes autonomes : un routeur frontière est désigné Forwarder de système autonome, et reçoit tous les paquets *multicast* circulant au sein du système autonome (et non plus seulement au sein d'une aire). La différence est importante, car la concentration de trafic est encore plus prononcée et les routeurs de ce type sont encore plus sollicités que les autres. Il est donc conseillé de dédier un routeur frontière de système autonome à cette seule tâche et de le relier à un réseau à haut débit.

Comme il a été mentionné au chapitre 11, OSPF ne permet pas de router les paquets entre AS. Il faut lui adjoindre un protocole de routage extérieur, tel que BGP4. Dans le cadre d'un réseau multicast, c'est DVMRP qui se charge de cette tâche au niveau des routeurs frontières.

Figure 13-11.
*Routage DVMRP
entre systèmes
autonomes.*



La configuration suivante doit donc être effectuée sur nos routeurs frontières RF1 et RF3 (ici, l'exemple concerne RF1) :

```
#Active la diffusion des routes MOSPF dans DVMRP
add -dvmrp mospf 220.10.0.0/16 aggregate 1
#Active le routage d'ospf vers dvmrp
setdefault -dvmrp policycontrol = mospf

#Active la diffusion des routes DVMRP dans ospf
add -mospf dvmrp 220.20.0.0/16 aggregate 5
#Active le routage de dvmrp vers ospf
setdefault -mospf policycontrol = dvmrp
```

L'option "Aggregate" permet de diffuser une seule route pour les routes comprises entre 220.10.0.0 et 220.10.255.0 au lieu de diffuser 256 routes.

Le système permet de mêler des routeurs OSPF et MOSPF. Cependant, aucune machine située derrière un routeur OSPF ne recevra les messages multicast. Il doit donc y avoir continuité des routeurs MOSPF tout au long des routes menant vers les machines multicast. À l'inverse de DVMRP, le tunneling n'est, en effet, pas supporté.

De plus, il faut toujours que le routeur désigné sur un réseau local soit un routeur MOSPF, sinon les requêtes IGMP n'alimenteront pas le protocole de routage multicast. Pour cela, il faut assigner une priorité supérieure aux routeurs OSPF également configurés en MOSPF.

Le routage à l'aide de PIM

Comme cela a été montré aux paragraphes précédents, le protocole MOSPF — et encore moins DVMRP — n'est pas adapté pour de grands réseaux, et notamment l'Internet (plusieurs dizaines de milliers d'aires et de systèmes autonomes).

Pour résoudre ce problème, l'IETF a défini le protocole PIM (*Protocol Independent Multicast*) qui, comme son nom le laisse penser, est indépendant du protocole de routage *unicast* utilisé. PIM fonctionne, en effet, avec RIP, OSPF, BGP, et même DVMRP.

PIM utilise deux modes de fonctionnement :

- *dense* (densité élevée), qui est adapté à des réseaux à haut débit et à des situations où les membres de groupes sont géographiquement proches ;
- *sparse* (clairsemé), qui correspond aux cas de figures où les réseaux ont de plus faibles débits et où les membres de groupes sont très dispersés.

Un routeur PIM bascule d'un mode à l'autre indépendamment de chaque groupe.

Le mode dense (PIM-DM) fonctionne de manière identique à DVMRP (algorithme *reverse path flooding*) sans toutefois utiliser un protocole de routage multicast dédié ; il utilise les protocoles de routage unicast en place (RIP, OSPF, etc.), d'où le nom de "*Protocol Independent*". PIM-DM ne fait pas encore l'objet d'une RFC.

Le mode *sparse* (PIM-SM) nécessite de convenir d'un routeur qui tienne lieu de point de rendez-vous pour chaque groupe. Un routeur peut être le point de rendez-vous pour plusieurs groupes. Il peut exister plusieurs points de rendez-vous pour un groupe, mais un seul doit être actif à un moment donné.

Principe de PIM-SM

Pour un couple réseau source/groupe multicast, un routeur PIM-SM bascule entre deux algorithmes de routage : un dont le chemin passe par le point de rendez-vous, et un second qui bénéficie d'un chemin direct entre la source et la destination. Pour le premier, le calcul de la route s'établit à partir d'un arbre partagé dont la racine est le point de rendez-vous (*Rendez-vous Point shared tree*). Pour le second, ce calcul se fonde sur un arbre du plus court chemin dont la racine est la source d'émission du paquet multicast (*source shortest path tree*).

Si plusieurs routeurs sont connectés au même réseau local, celui dont l'adresse IP est la plus élevée sera élu routeur désigné et aura à charge d'envoyer et de recevoir les messages Join/Prune. Pour découvrir ses voisins, un routeur émet périodiquement un message Hello dont la périodicité par défaut est de 30 secondes :

```
ip pim query-interval 30
```

Tous les routeurs désignés dont les membres actifs appartiennent à un même groupe s'enregistrent auprès du même routeur, appelé point de rendez-vous (RP). Les routeurs intermédiaires enregistrent également cette information et font de même auprès du RP. Ce dernier peut donc calculer un arbre de routage pour un couple adresse réseau source/adresse multicast de groupe. En définitive, le RP connaît tous les réseaux sources susceptibles d'émettre et de recevoir des paquets multicast au sein d'un groupe donné.

Pour terminer cette introduction, plusieurs routeurs peuvent être candidats à l'élection du point de rendez-vous. L'un d'eux est élu routeur bootstrap ; il est chargé de désigner le RP actif pour chaque groupe. Il en résulte de nouveaux messages d'annonce entre les RP.

Principe du calcul des routes

Un routeur PIM-SM ayant enregistré (*via* IGMP) des membres d'un ou plusieurs groupes envoie périodiquement un message Join/Prune (joindre/élaguer) au point de rendez-vous RP. Chaque routeur chargé de transporter ce message se trouvant sur le chemin enregistre les mêmes informations : le groupe, la source (l'adresse du réseau sur lequel ont été identifiés les membres du groupe) et l'interface d'entrée du message Join/Prune.

Un message Join/Prune contient, pour chaque adresse de groupe (multicast), la liste des adresses sources incluses dans l'arbre de routage (*joined*, qui ont rejoint l'arbre) et de celles qui ne le sont pas (*pruned*, élaguées de l'arbre).

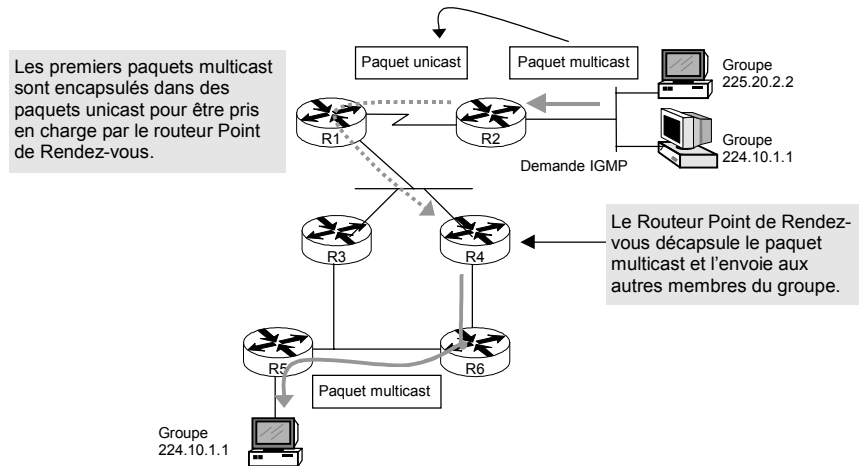
Un routeur émet un message Join/Prune vers un routeur situé en amont (en direction de la racine de l'arbre). Il indique par là qu'il fait partie de l'arbre de routage pour les paquets *multicast* du groupe G émis par la source S.

Principe du routage

Le premier paquet multicast émis par un PC est envoyé dans un paquet multicast IP qui est pris en charge par le routeur PIM-SM désigné. Celui-ci encapsule ce paquet dans un paquet unicast à destination du RP. Le RP décapsule ce paquet et l'envoie tel quel (donc sous forme multicast) aux membres du groupe (les routeurs qui se sont déclarés).

Figure 13-13.

Principe
du routage
PIM-SM.



Tous les messages passent donc par le RP, ce qui n'est pas forcément une route optimale.

Si le trafic multicast persiste, un routeur (la source, la destination ou le RP) peut alors décider que le flux emprunte un chemin direct entre la source et la destination. Les routeurs concernés basculent alors d'un arbre de routage dont la racine est le RP vers un arbre du plus court chemin dont la racine est la source.

Dans notre réseau, le routeur R1 crée la table de routage suivante :

```
Router4# show ip pim interface
Address      Interface    Mode    Neighbor Count    Query Interval    DR
220.20.20.6  Ethernet0    sparse   2                30                220.20.20.6
197.10.10.0  serial0      sparse   1                10                0.0.0.0
```

La colonne *Address* indique l'adresse du routeur vers lequel envoyer les paquets à router *via* l'interface spécifiée dans la colonne *Interface*.

La colonne *Neighbor count* indique le nombre de voisins découverts sur cette interface.

La colonne *DR* indique l'adresse du routeur désigné.

Routage sur les liaisons WAN

Sur Ethernet, le paquet IP multicast est envoyé dans une trame Ethernet multicast. Mais, sur les réseaux qui n'offrent pas cette fonction (X.25, Frame Relay, ATM, etc.), seules des trames unicast peuvent être envoyées.

LE POINT SUR PIM-SM (RFC 2362)

Un routeur source qui reçoit des rapports IGMP indiquant que des membres du groupe G sont actifs ajoute une route (*, G) vers le routeur **point de rendez-vous** pour le groupe ; il lui envoie alors périodiquement un message **Join/Prune** (joindre/élaguer). Tous les routeurs situés sur le chemin qui relie ces deux routeurs ajoutent la même route (*,G) en notant l'adresse du réseau source et l'interface d'entrée, puis envoient un message Join/Prune au routeur point de rendez-vous. Ce dernier connaît donc le chemin le reliant au routeur source qui déclare que des membres du groupe G sont actifs. Les routeurs ont donc rejoint un **arbre partagé** qui a comme racine le routeur point de rendez-vous (*RP shared tree*).

Un routeur source ayant pris en charge un paquet *multicast* l'encapsule dans un paquet *unicast*, puis l'envoie au routeur point de rendez-vous, qui le décapsule et le renvoie vers tous les routeurs sources déclarés (qui sont ici des routeurs destinations) et ayant rejoint l'arbre partagé.

Au bout d'un moment, un routeur (source, destination ou point de rendez-vous) peut décider de quitter l'arbre partagé afin que les paquets *multicast* soient échangés directement sans passer par le routeur point de rendez-vous. Il envoie alors un message Join/Prune au point de rendez-vous pour quitter l'arbre partagé, et un autre à la source pour rejoindre l'**arbre du plus court chemin** (SPT, *Shortest Path Tree* — dont la racine est la source) pour le couple routeur source S/groupe G. Tous les routeurs situés sur le chemin routeur source-routeur destination ajoutent la route (S,G) dans leur table de routage.

Le basculement vers le SPT s'opère lorsque le flux multicast est suffisamment important et de longue durée (en fait, au bout de quelques secondes).

Les messages Join/Prune contiennent, pour chaque groupe, la liste des sources (routeurs ou réseau) appartenant à l'arbre et celles qui en sont élaguées.

4 bits		4 bits		8 bits		16 bits	
Version	Type	Longueur adresse		Checksum			
Adresse du routeur upstream							
Réservé		Nombre de groupes		Durée de vie de l'état Join/Prune			
1 ^{ère} adresse multicast de groupe							
Nombre de sources dans l'arbre				Nombre de sources élaguées			
Adresse de la source 1 ayant rejoint l'arbre (Join)							
...							
Adresse de la source 1 élaguée (Prune)							
...							
2 ^{ème} adresse multicast de groupe							
...							

Le champ "Nombre de groupes" indique le nombre de groupes décrits dans le message (notés 1, 2, etc.).

Le champ "Adresse de la source" indique que le routeur transmettra (ou non, s'il est dans l'état Prune) un paquet multicast issu de cette source s'il provient de l'interface par laquelle il envoie ce message Join/Prune. Une adresse source peut être celle d'un routeur, d'un réseau ou d'un agrégat de réseaux.

Les routeurs diffusent les paquets multicast uniquement vers les interfaces par lesquelles sont entrés des messages Join/Prune. Par ce biais, chaque routeur connaît la topologie du réseau et calcule un arbre de routage (soit partagé, soit du plus court chemin) pour chaque couple Source/Groupe. Les routeurs s'échangent, par ailleurs, des messages Hello pour découvrir leurs voisins.

Les messages PIM sont envoyés dans des paquets *unicast* (*Register* et *Register-stop*) et *multicast* 224.0.0.13 (Hello, etc.) en utilisant le protocole n° 103 au-dessus d'IP.

L'activation du mode NBMA (*NonBroadcast MultiAccess*) permet au routeur de garder une trace des adresses IP qui émettent les messages PIM *Join* arrivant par l'intermédiaire de l'interface WAN. Les paquets multicast devant ainsi être diffusés vers cette interface seront dupliqués en autant de paquets encapsulés dans des trames unicast.

Cela est notamment le cas sur les interfaces série de nos routeurs R1 et R2 :

```
int s 0
ip pim nbma-mode
```

Il est à noter que le mode NBMA consomme de la mémoire puisqu'il garde trace de toutes les adresses IP sources arrivant par l'intermédiaire de l'interface WAN.

De même, le mode *source shortest path tree* requiert davantage de mémoire que le mode *RP shared tree*. En revanche, il réduit la charge réseau et optimise le chemin (il n'est plus besoin de passer par le RP).

Les routeurs Cisco basculent du mode *shared* vers le mode *shortest path* dès le premier paquet émis. Cependant, il est possible de n'activer le basculement qu'à partir d'un certain débit constaté, en d'autres termes, lorsque cela en vaut la peine.

```
ip pim spt-threshold 64 group-list 1
```

La commande précédente indique que le basculement s'opérera lorsque le flux multicast aura atteint 64 Kbit/s pour les groupes figurant dans l'*access-list* 1.

Quel protocole choisir ?

Pour l'enregistrement des membres d'un groupe auprès d'un routeur, IGMP est incontournable. Les informations collectées sont ensuite utilisées par les protocoles de routage multicast.

Trois protocoles de routage multicast sont proposés : trois solutions, trois approches différentes.

Critère	DVMRP	MOSPF	PIM
Mise à jour des tables de routage	Diffusion périodique de l'intégralité des tables	Diffusion des modifications par messages d'annonce	Centralisée au niveau du point de rendez-vous (PR)
Diffusion du premier paquet	Inondation du premier paquet	Routé directement vers les membres du groupe	Encapsulé dans un paquet unicast vers le PR, puis re-diffusé par le PR
Calcul du meilleur chemin (arbre SPT) pour le couple Source/Groupe	Par élagage des routeurs ne désirant pas recevoir le paquet	À la demande, par tous les routeurs recevant le premier paquet	Par remontée de messages Join/Prune de proche en proche jusqu'à la source
Protocole de routage	DVMRP	MOSPF s'appuie sur OSPF	S'appuie sur les protocoles existants (RIP, OSPF, etc.)
Tunneling entre routeurs non multicast	Oui	Non : les routeurs MOSPF doivent être contigus	Oui

DVMRP est le protocole le plus ancien et le moins performant (il présente les mêmes défauts que RIP). L'inondation régulière des paquets multicast ainsi que la diffusion périodique des tables de routage ne font pas de DVMRP un protocole adapté aux grands réseaux.

MOSPF est le protocole le plus performant au sein d'une aire ; il est relativement performant entre aires. Entre systèmes autonomes, l'usage de DVMRP est requis pour le multicast (en plus de BGP pour l'unicast), ce qui peut présenter quelques inconvénients.

PIM est adapté aux grands réseaux et notamment l'Internet, là où les membres d'un même groupe sont très dispersés. Dans un petit réseau, cet avantage se transforme en inconvénient, car le routeur de point de rendez-vous n'est pas obligatoirement situé sur le meilleur chemin (à moins que le réseau soit très centralisé).

D'autres considérations, non techniques, entrent en jeu :

- Les RFC de DVMRP (datée de 1988) et de PIM (datée de 1998) n'en sont qu'au stade expérimental, alors que celle de MOSPF (datée de 1994) en est à l'état de proposition de standard.
- Cependant, Cisco, le principal fournisseur de routeurs sur l'Internet, ne prend en charge que PIM, et pour cause : les ingénieurs de la société ont participé à son élaboration mais pas à celle de MOSPF.
- En outre, DVMRP est le plus répandu au sein de MBONE (*Multicast Backbone*, une portion de l'Internet expérimentant le multicast). Son utilisation est donc requise, au moins sur le routeur de frontière utilisé en interconnexion avec ce réseau.

En conclusion :

- Si vos routeurs sont de marque Cisco, PIM est le seul choix possible.
- Si votre réseau fonctionne avec OSPF et ne comprend pas de routeurs Cisco, le choix de MOSPF va de soi.
- Utilisez DVMRP entre les systèmes autonomes et pour vous raccorder à MBONE.

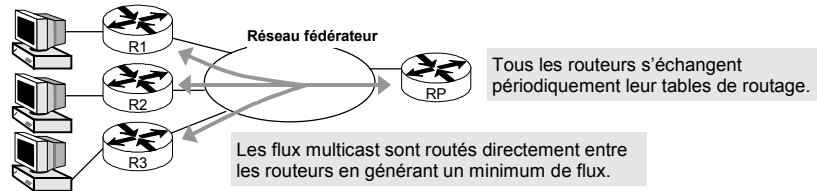
L'étape suivante consiste à définir l'architecture la mieux adaptée au protocole choisi, afin de limiter le trafic généré par les routeurs.

Si plusieurs routeurs sont présents sur le même réseau local, il faut déterminer quel est le meilleur candidat à l'élection du demandeur. Le processus d'élection dépend du protocole de routage activé. Pour les routeurs DVMRP, le routeur demandeur IGMP est celui qui a la plus petite adresse IP ou celui dont la métrique est la plus faible. Pour les routeurs MOSPF, il s'agit du routeur désigné OSPF (celui dont la priorité est la plus basse), tandis que pour PIM, il s'agit de celui dont l'adresse IP est la plus haute.

Architecture adaptée au protocole DVMRP

Il existe peu de solutions pour optimiser les flux DVMRP, si ce n'est l'emploi d'un réseau centralisé autour d'un réseau fédérateur.

Figure 13-14.
*Optimisation
des flux DVMRP.*

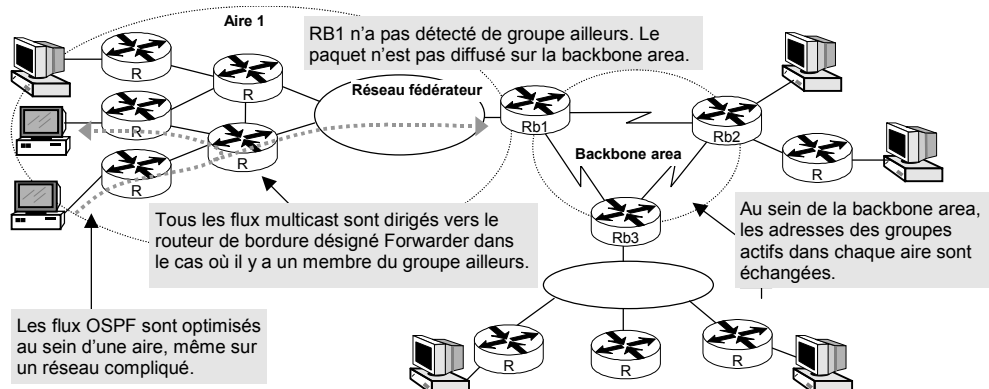


On le voit, DVMRP est adapté à un réseau localisé sur un site ou un campus avec un réseau fédérateur à haut débit.

Architecture adaptée au protocole MOSPF

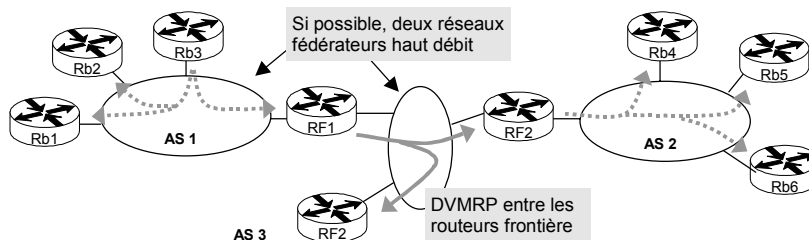
Le routage au sein des aires est bien contrôlé par MOSPF. Entre les aires, le routeur de bordure désigné forwarder reçoit tous les flux multicast ; il doit donc être positionné sur un réseau fédérateur à haut débit.

Figure 13-15.
*Optimisation
des flux MOSPF.*



MOSPF est adapté aux grands réseaux mais limité à un seul système autonome. En effet, la concentration des flux entre AS est encore augmentée, à l'image de la concentration sur la backbone area.

Figure 13-16.
*Optimisation des flux
entre systèmes autonomes.*



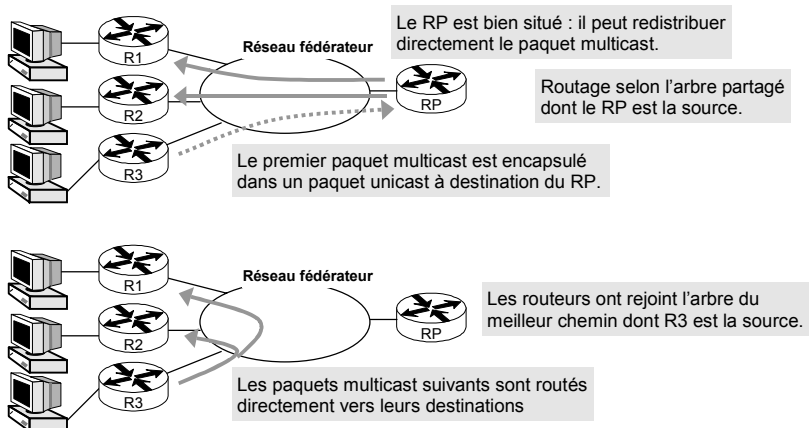
Entre systèmes autonomes, DVMRP doit être configuré en point à point, ou en multipoint sur un réseau fédérateur à haut débit, et ce afin d'optimiser le flux multicast.

Afin de limiter la taille des tables de routage, les routes redistribuées entre DVMRP et MOSPF doivent être agrégées. Cela implique que les adresses réseau soient contiguës au sein d'un système autonome (voir chapitre 11).

Architecture adaptée au protocole PIM

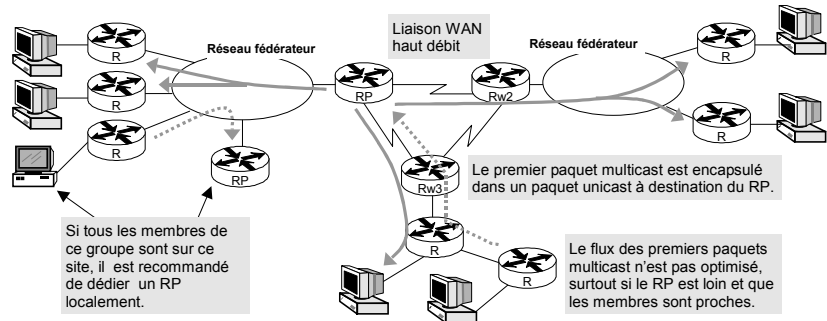
Le routeur point de rendez-vous de PIM doit être situé au centre du réseau, afin d'optimiser la route des premiers paquets multicast.

Figure 13-17.
*Optimisation
des flux PIM.*



Si les membres d'un groupe sont répartis entre plusieurs sites, la meilleure position du RP se situe sur un routeur WAN.

Figure 13-18.
Flux PIM
sur un réseau intersite.



Il est conseillé de paramétrer les routeurs PIM pour basculer rapidement vers l'arbre du meilleur chemin, afin de ne plus passer inutilement par le RP. Les routeurs Cisco basculent dès le premier paquet : la commande `ip pim spt-threshold` ne doit donc pas être utilisée. Le revers de la médaille est un basculement inutile si le flux multicast est sporadique, ce qui entraîne une plus grande consommation de bande passante réseau (pour les messages Join/Prune) et de ressources CPU sur tous les routeurs qui calculent le nouvel arbre du meilleur chemin. En outre, ce dernier utilise beaucoup de ressources mémoire, d'autant plus qu'un arbre par couple adresse source/adresse de groupe est nécessaire.

Par ailleurs, des messages Bootstrap circulent entre les routeurs point de rendez-vous, ajoutant encore des flux de gestion propres à PIM.

Tout est donc affaire de dosage en fonction du nombre d'utilisateurs et du type de trafic. Du fait de la complexité du protocole PIM, la conception de l'architecture réseau n'en est que plus difficile.

Contrôler la diffusion sur son réseau

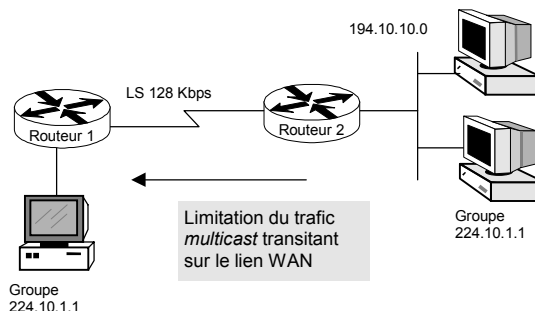
Les flux *multicast*, tels que la vidéo, peuvent générer un volume de données important. Afin de contrôler ces flux, il est possible de filtrer les groupes multicast sur chaque interface. Cela évite également de voir fleurir des groupes un peu partout sans que l'administrateur n'en soit tenu informé.

```
access-list 90 225.20.2.2 0.0.0.0
access-list 90 224.10.1.1 0.0.0.0
interface ethernet 0
ip igmp access-group 90
```

L'access-list 90 contient la liste des adresses de groupe qui seront autorisées à être diffusées sur l'interface Ethernet 0.

Il est également intéressant de contrôler le débit généré par les membres d'un groupe, pour ne pas pénaliser les autres applications et, également, ne pas inonder votre réseau.

Figure 13-19.
Contrôle du trafic multicast.



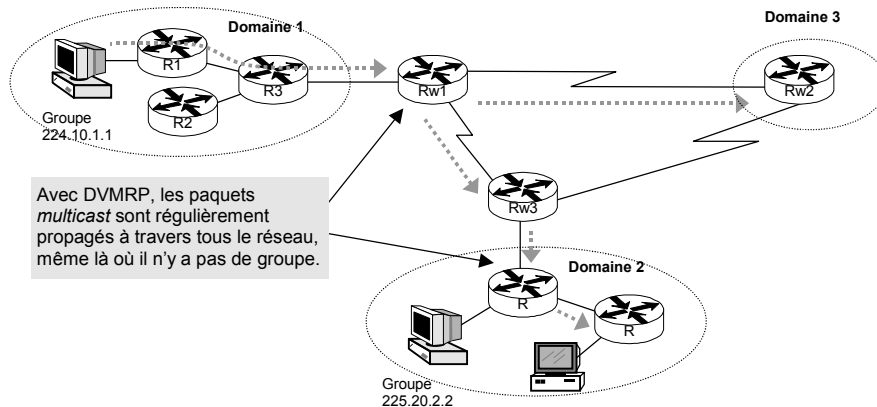
Dans cet exemple, les membres indiqués dans l'access-list 90 et qui émettent au sein des groupes précisés dans l'access-list 80 sont limités à un débit de 64 Kbit/s sur la liaison WAN. Cette limite peut être positionnée en entrée (*in*) ou en sortie (*out*) de l'interface.

```
interface serial 1
ip multicast rate-limit out group-list 80 source-list 90 64
access-list 80 permit 0.0.0.0 255.255.255.255
access-list 90 permit 194.10.10.0 0.255.255.255
ip multicast rate-limit in group-list 80 source-list 90 64
```

Dans l'exemple précédent, tous les multicast issus du réseau 194.10.10.0 seront limités à 64 Kbit/s, laissant une bande passante de 64 Kbit/s disponible pour les autres applications.

Reprenons l'exemple de notre réseau intersite.

Figure 13-20.
Contrôler la diffusion.

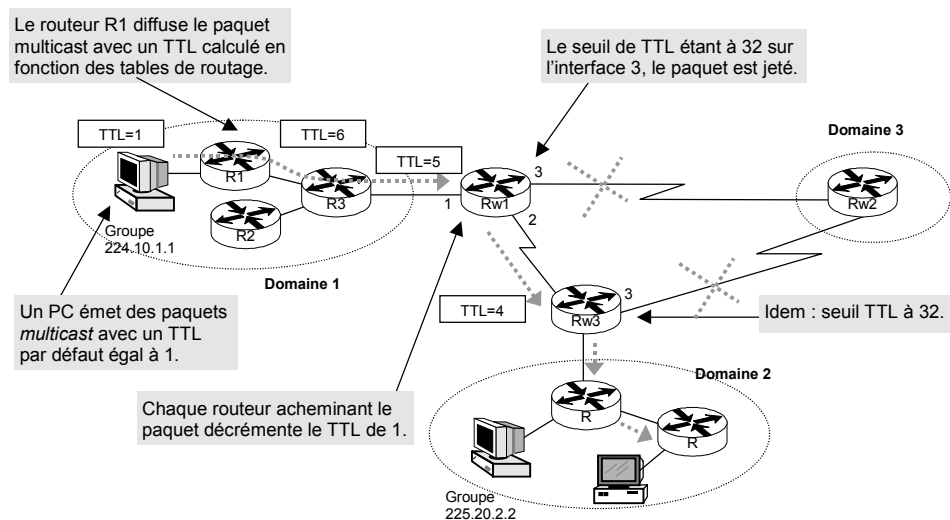


Si DVMRP est utilisé, tous les premiers paquets multicast émis par une source pour un groupe donné sont diffusés à travers l'ensemble du réseau (tant qu'il y a des routeurs DVMRP). Comme on l'a vu, DVMRP est obligé de répéter cette procédure régulièrement afin de tenir compte des changements de topologie et d'appartenance à un groupe.

Cela est également le cas de MOSPF, dans la *backbone area*, et de PIM, au niveau des réseaux auxquels est connecté le routeur point de rendez-vous.

Les routeurs permettent de limiter la propagation de ces paquets en contrôlant leur TTL (*Time To Live*). Ce mécanisme impose qu'un paquet IP ne soit pas routé lorsque son TTL est égal à 1. Le mécanisme proposé permet de définir un seuil supérieur en dessous duquel le paquet ne sera pas routé.

Figure 13-21.
*Limitation de la diffusion
des paquets.*



Seuls les paquets dont le TTL est supérieur au seuil de 32 seront routés par les routeurs Rw1 et Rw3 sur leur interface série 3. La commande Cisco est la suivante :

```
int s3
ip multicast ttl-threshold 32
```

La valeur par défaut est 0

En positionnant différents seuils à certains points stratégiques du réseau, il est possible de définir des domaines de diffusion. On peut ainsi limiter la portée du point de rendez-vous PIM ou la propagation des paquets *via* des routeurs DVMRP.

Par exemple, MBONE utilise la convention suivante :

TTL	Périmètre
0	Restreint à la machine (paquet boucle en local)
1	Restreint au même réseau local
32	Site
64	Région
128	Continent
255	Pas de restriction de diffusion

Des mécanismes propres à IGMP sont également prévus pour limiter le trafic local :

- Les rapports sont envoyés avec un TTL = 1, ce qui permet de supprimer le paquet rapidement : le rapport est censé être adressé uniquement au routeur du réseau local.
- Les rapports sont envoyés un à un pour chaque groupe d'appartenance, avec un intervalle de temps aléatoire.
- Au niveau du routeur, les temporisateurs (*timers*) sont armés indépendamment pour chaque groupe identifié.

Les groupes de paquets compris entre 224.0.0.0. et 224.0.0.255 ne sont jamais routés : de tels paquets transitent d'une machine vers un routeur, d'un routeur vers un autre routeur, et ce quelle que soit la valeur du TTL (en principe égale à 1).

La qualité de service sur IP

Par défaut, un réseau IP se contente d'acheminer les paquets au mieux de ses possibilités, et sans distinction. Tant que la bande passante (c'est-à-dire le débit) est suffisante, il n'y a pas de problème. Mais, en cas de saturation, les routeurs sont obligés de rejeter des paquets, invitant tous les émetteurs à réduire leur flux. En conséquence, l'utilisateur constate une dégradation des performances du réseau.

La notion de qualité de service (QoS, *Quality of Service*) introduit la possibilité de partager le plus équitablement possible une ressource devenant de plus en plus rare, car partagée par un grand nombre de flux applicatifs qui peuvent interférer les uns avec les autres. Elle introduit également la possibilité de déterminer différents niveaux de service en fonction de la nature de ce flux (une visioconférence, un transfert de fichier, etc.).

Au chapitre 10, vous avez sans doute remarqué que la gestion de la qualité de service est déjà prise en compte par des protocoles de niveau 2, tels qu'ATM et Frame Relay. Alors pourquoi gérer la QoS sur IP ? Parce que, lorsqu'une application génère des flux sur un réseau Ethernet, qui traversent ensuite un réseau ATM ou Frame Relay pour arriver sur un autre réseau local, le seul dénominateur commun est IP.

Dans ce chapitre, vous apprendrez ainsi :

- comment améliorer les performances de votre réseau ;
- les différents moyens permettant de gérer la qualité de service ;
- quelle politique de qualité de service choisir ;
- comment configurer votre réseau pour gérer cette qualité de service.

Améliorer les performances du réseau

Notre réseau est de plus en plus sollicité par de nouvelles applications aux besoins très divers : connexions Telnet pour se connecter à une machine Unix, transferts de fichiers, bases de données en mode client-serveur et, maintenant, flux audio et vidéo.

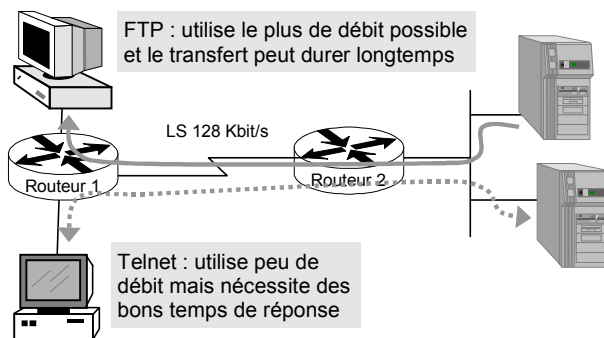
La charge du réseau augmente, et les flux générés se perturbent mutuellement : quel utilisateur n'a pas expérimenté des temps de réponse très longs alors que son voisin a lancé le transfert d'un gros fichier ?

Affecter des priorités sur les files d'attente

Le moyen le plus rudimentaire d'améliorer la qualité du service rendu par votre réseau est de définir manuellement, et sur tous vos routeurs, des priorités. Cette approche consiste surtout à améliorer les performances du réseau en favorisant des applications au détriment d'autres. La qualité est améliorée mais non garantie.

Le cas le plus couramment rencontré est celui d'une liaison WAN qui doit véhiculer des connexions Telnet (de type conversationnel) et des transferts de fichiers FTP. Les temps de réponse Telnet se dégradent dès qu'un transfert FTP est lancé. Aucune qualité de service n'étant gérée, la bande passante du réseau est, en effet, accaparée par celui qui en consomme le plus, cela bien sûr au détriment des autres.

Afin d'éviter que les applications Telnet ne soient gênées par un transfert de fichiers, le moyen le plus simple est de définir des priorités sur la base des ports TCP qui identifient les applications (23 pour Telnet et 20 pour le canal de données FTP).



En classant les paquets Telnet en priorité haute, les paquets IP comportant le port TCP 23 seront placés dans la file d'attente à priorité haute, tandis que les paquets FTP seront stockés dans la file d'attente à priorité basse. Nos routeurs, qui sont de marque Cisco, gèrent ainsi quatre files d'attente correspondant à quatre priorités : haute, moyenne, normale et basse (*high, medium, normal, low*).

Il existe un jeu de files d'attente par interface réseau. La priorité doit donc être activée sur les interfaces série qui émettent les flux FTP perturbateurs. Étant donné que le transfert de fichiers peut avoir lieu dans les deux sens, la même configuration doit être appliquée aux deux routeurs :

```
priority-list 1 protocol ip high tcp 23
priority-list 1 protocol ip low tcp 20
priority-list 1 protocol ip low tcp 21
int s 0
priority-group 1
```

Telnet

FTP

En positionnant une priorité basse pour les flux FTP, les transferts de fichiers prendront plus de temps mais ne perturberont pas les sessions Telnet. Lorsque plusieurs paquets sont en file d'attente, le routeur enverra de préférence davantage de paquets Telnet que de paquets FTP.



Afin d'offrir le même comportement sur l'ensemble du réseau, il faut configurer de manière identique tous les routeurs. De même, les interfaces Ethernet pourraient être traitées de manière identique si les réseaux locaux étaient chargés. En fonctionnement normal, le débit de 10 Mbit/s suffit, en effet, à absorber les flux FTP et Telnet. Ce n'est qu'en cas de charge que la perturbation se manifeste et que l'activation des priorités permet de conserver le même niveau de service que celui obtenu en fonctionnement non chargé.

Il est intéressant de spécifier plus en détail un flux en utilisant une access-list :

```
priority-list 1 protocol ip low list 10
access-list 10 permit 10.0.0.1 0.255.255.255
```

Par défaut, tous les autres paquets seront traités de la manière suivante :

```
priority-list 1 default medium
```

Il est également possible de contrôler indirectement le débit en définissant le nombre maximal de paquets pouvant être en attente :

```
priority-list 1 queue-limit 20 40 60 80
```

Les quatre chiffres indiquent le nombre maximal de paquets pouvant être stockés dans les files d'attente haute, moyenne, normale et basse (les valeurs indiquées sont celles par défaut). Si la file d'attente est pleine, les paquets en excès sont rejetés, et un message ICMP *source-quench* est envoyé à l'émetteur pour lui indiquer de ralentir le flux.

Cette stratégie a des limites, car elle est à double tranchant lorsque le flux de paquets est important : si la file d'attente est de trop grandes dimensions, les paquets s'accumulent, ce qui a pour conséquence d'augmenter le temps de réponse global. Si, en revanche, elle est de trop

petite taille, des paquets peuvent être perdus, notamment entre deux réseaux de débits différents (depuis un réseau Ethernet à 10 Mbit/s vers un lien série, par exemple) : le routeur rejette tout paquet entrant tant que ses files d'attente sont saturées.

Dans notre exemple, en abaissant ainsi à 60 le nombre de paquets dans la file d'attente basse, le débit maximal du flux FTP sera encore abaissé.

Agir sur les files d'attente

L'étape suivante vers une gestion de la qualité de service consiste à agir sur le comportement des files d'attente des routeurs.

LE RÔLE DES FILES D'ATTENTE

Lorsqu'il arrive plus de paquets qu'il n'en sort du routeur, celui-ci les garde en mémoire en attendant que les plus anciens soient envoyés. Si d'autres paquets continuent d'arriver, la file d'attente se sature. Le débordement de la file d'attente se traduit par le rejet des paquets continuant d'arriver (le routeur les ignore).

Un algorithme de traitement des files a donc deux rôles essentiels :

- **traiter** en priorité tel ou tel paquet en cas de **congestion** ;
- **rejeter** en priorité tel ou tel paquet en cas de **saturation** de la file d'attente.

La priorité de traitement d'un paquet dépend de paramètres de qualité de service qui peuvent être prédéfinis dans l'algorithme, soit définis statiquement dans le routeur, soit définis dynamiquement par l'application.

Les paquets rejetés signifient pour TCP qu'il doit réduire sa fenêtre d'émission et donc son flux. Si le même paquet est rejeté plusieurs fois, l'utilisateur attend, et si l'attente se prolonge, la session risque d'être interrompue à l'expiration d'un timeout. Le choix des paquets à rejeter dépend de l'**algorithme** choisi pour traiter la file d'attente.

L'algorithme FIFO – Un fonctionnement simple

Le moyen le plus simple est de gérer les files d'attente sur le mode FIFO (*First In, First Out*) : le routeur traite les paquets au fil de l'eau dans leur ordre d'arrivée, au mieux de ses capacités. C'est le principe du *best effort*, comportement par défaut des routeurs. Aucun paramétrage n'est possible, si ce n'est de définir des priorités ce qui a pour effet de créer une file d'attente par niveau de priorité.

Gérer les congestions

Une méthode plus efficace est de traiter les files d'attente à l'aide de l'algorithme WFQ (*Weighted Fair Queueing*). Celui-ci identifie dynamiquement les flux et veille à ce que les applications générant peu de trafic ne soient pas perturbées par celles générant beaucoup de données.

Le principe de l'algorithme étant figé, la possibilité de paramétrage est donc limitée au seuil de rejet et au nombre de files d'attente :

```
int s 0
bandwidth 512
fair-queue 64 256 0 (valeurs par défaut)
```

Dans l'ordre, les paramètres indiquent :

- le seuil au-delà duquel les paquets vont commencer à être rejetés (c'est-à-dire la taille de la file d'attente moins une petite marge), ici 64 paquets par file ;
- le nombre de files d'attente pouvant être créées dynamiquement pour les flux sans qualité de service (de type *best effort*), ici 256 files pour traiter simultanément 256 flux ;
- le nombre de files d'attente pouvant être réservées par RSVP (voir plus loin), dans notre cas aucune.

Les paquets qui disposent des mêmes adresses IP source et destination, des mêmes ports source et destination et du même champ TOS (voir plus loin) correspondent à un même flux. Une fois identifiés, tous les paquets du même flux sont placés dans la même file d'attente.

La commande “bandwidth”, qui indique le débit du lien réseau en Kbit/s, permet à l'algorithme de définir le nombre nécessaire de tampons d'émission associés à l'interface (généralement quelques-uns).

Le mot “*weighted*” dans WFQ indique que l'algorithme prend en considération la priorité indiquée dans le champ “*IP precedence*” du paquet (voir le paragraphe suivant à ce sujet).

Prévenir les congestions

Alors que l'algorithme WFQ permet de gérer les situations de congestion, l'algorithme WRED (*Weighted Random Earle Detection*) permet de les prévenir. Dès qu'une congestion est détectée, l'algorithme rejette des paquets, ce qui contraint l'émetteur à ralentir son flux :

```
int s 0
bandwidth 512
random-detect
```

Tout comme WFQ, WRED permet de prendre en compte la priorité du paquet définie dans le champ “*IP precedence*”.

Réguler le trafic

Une autre alternative est de réguler le flux selon le principe du *traffic shaping* : le but de cet algorithme est de donner un caractère prévisible aux flux en convertissant un flux erratique en un flux à peu près constant.

Sur nos routeurs, l'algorithme peut être appliqué à l'ensemble du trafic, ou à une portion de celui-ci définie par une *access-list* :

```
access-list 101 permit ftp any any
int s 0
```

```
traffic-shape group 101 128000 64000 32000
```

```
int e1
```

```
traffic-shape rate 5000000 10000000 20000000
```

Taille du *burst* en bits

Définitions identiques à celles du Frame Relay décrit au chapitre 10.

Burst en excès = nombre de bits pouvant dépasser le débit moyen.

La première commande indique que le débit moyen est limité à 128 Kbit/s, mais qu'un dépassement de 32 kilo-bits est autorisé pendant une demi-seconde (64/128). Cela signifie que le débit peut atteindre 192 Kbit/s ($128 + 32/0,5$) pendant une demi-seconde. Cette régulation s'applique aux flux FTP tel qu'indiqué dans l'*access-list* 101.

La deuxième commande est appliquée à une interface Ethernet afin de limiter le débit vers ce réseau à 5 Mbit/s en autorisant une pointe à 6 Mbit/s ($5 + 2/(10/5)$) pendant 2 secondes (10/5).

LA GESTION DES FILES D'ATTENTE

Comme dans la vie de tous les jours, une file d'attente se forme lorsque le flux entrant est plus important que le rythme de sortie des paquets, par exemple depuis une interface Ethernet à 10 Mbit/s vers une liaison série à 512 Kbit/s. Le routeur peut alors adopter différentes stratégies pour traiter les paquets en attente. Même si tout se passe en une fraction de seconde, le choix de l'une ou de l'autre d'entre elles peut avoir une grande influence sur le comportement général du réseau.

Le mode le plus simple est de type **FIFO** (*First In First Out*) qui consiste à traiter les paquets dans leur ordre d'arrivée. C'est celui qui consomme le moins de CPU et qui engendre le moins de latence pour les paquets. En cas de congestion, il se révèle, en revanche, moins performant car il n'y a aucune régulation du trafic.

Le principe du **WFQ** (*Weighted Fair Queueing*) repose quant à lui sur un contrôle de flux dynamique en fonction de discriminants propres à chaque protocole (le DLCI Frame Relay, les adresses IP et le port TCP, une priorité, etc.), le principe étant de privilégier le trafic à faible volume sur celui à fort volume. Pour cela, chaque flux est identifié (à partir des adresses IP source et destination et des port TCP/UDP source et destination), et son débit est mesuré.

Le principe de l'algorithme **RED** (*Random Early Detection*) consiste à prévenir les congestions. Lorsque la file d'attente commence à être saturée, des paquets correspondant à des flux sélectionnés aléatoirement sont rejetés. Les variantes **WRED** (*Weighted RED*) et **ERED** (*Enhanced RED*) permettent de sélectionner les flux, en fonction de priorités qui déterminent le rejet des paquets.

Le **Traffic shaping** est un mécanisme qui permet de réguler les flux de données, c'est-à-dire de fluidifier en sortie un trafic qui est erratique en entrée. La première implémentation, appelée **Leaky-bucket** (littéralement le seau percé), consiste à offrir un débit stable en sortie. La seconde variante, appelée **Token-bucket**, ne fluidifie pas le trafic erratique tant qu'il n'atteint pas un certain seuil (c'est-à-dire tant qu'il ne dépasse pas le nombre de jetons, un jeton équivalant à un certain nombre d'octets). Au-delà du seuil, tous les trafics sont fluidifiés. Selon les implémentations, le *Traffic shaping* utilise sa propre file d'attente ou peut opérer conjointement avec une file FIFO ou WFQ.

Quelle file d'attente choisir pour son réseau ?

Chacune des fonctions qui viennent d'être décrites présente des avantages et des inconvénients. Le problème est alors de savoir dans quelle situation et à quel endroit du réseau utiliser l'une ou l'autre de ces fonctionnalités.

Nous l'avons vu précédemment, une **file d'attente FIFO** ne permet pas de gérer des situations de congestion. En conséquence, une file d'attente de ce type ne peut être utilisée qu'avec des réseaux non saturés et offrant suffisamment de débit par rapport au trafic.

L'affectation de **priorités sur les files d'attente**, qui a été l'objet de notre première démarche, consiste en une programmation arbitraire, statique, et qui ne distingue qu'individuellement les types de paquets. Lorsqu'une congestion survient, les paquets prioritaires sont traités tant qu'il y en a dans la file d'attente correspondante. L'effet de bord qui en résulte est que les autres paquets restent bloqués dans la file d'attente, ce qui peut entraîner un ralentissement conséquent, voire la coupure des sessions correspondantes à ces paquets non prioritaires. C'est néanmoins la méthode d'affectation de priorité qui utilise le moins de CPU car l'algorithme est simple. L'utilisation de cette fonction sera donc limitée à des liens à bas débit sur lesquels le trafic est bien identifié.

La **file d'attente WFQ** identifie chaque flux, et traite tous les paquets d'un même flux de la même manière. En cas de congestion, l'algorithme traite équitablement tous les flux en privilégiant ceux à faible volume, mais pas au détriment de ceux à fort volume. Bien qu'efficace, cet algorithme complexe consomme de la ressource CPU, ce qui le destine plutôt à des liens à bas et moyen débits. Du fait qu'il prend en compte les priorités indiquées dans les paquets, on activera cet algorithme au sein du réseau et non dans sa périphérie.

La **file d'attente WRED** permet de prévenir les congestions en rejetant aléatoirement des paquets. À partir d'un certain seuil, le taux de rejet de paquets augmente à mesure que la file d'attente se remplit. Cet algorithme est de ce fait particulièrement bien adapté au protocole TCP qui est prévu pour réduire sa fenêtre d'émission en cas de perte de paquet (ce qui n'est pas le cas des protocoles IPX et AppleTalk).

À l'inverse de la file d'attente pour laquelle les priorités sont fixées dans le routeur, les algorithmes WFQ et WRED prennent en compte la priorité qui est indiquée dans le paquet IP.

Enfin, le **Traffic shaping** est un mécanisme qui permet de lisser le trafic et de fixer un débit à chaque type de flux. Il est donc bien adapté aux routeurs situés en bordure du réseau, c'est-à-dire sur les réseaux locaux, là où utilisateurs et serveurs émettent leurs données. Ce trafic erratique peut ainsi être régulé lorsqu'il entre dans le réseau longue distance, ce qui permet ensuite aux routeurs situés au sein de ce réseau d'utiliser les files d'attente WFQ et WRED.

La file d'attente de type...	... doit être utilisée
FIFO	Lorsque le réseau n'est pas saturé
Priorité sur file d'attente	Sur des liens à bas débit
WFQ	Au sein du réseau, sur des liens à bas et moyen débits (< 2 Mbit/s)
WRED	Au sein du réseau, sur des liens à haut débit
Traffic shaping	En entrée du réseau et pour adapter le flux entre des liaisons ayant des débits différents

Gérer la qualité de service

À la base, un réseau IP tel que l'Internet ne garantit pas que tous les paquets émis seront délivrés au destinataire. Il assure simplement que les paquets effectivement remis le seront sans erreur ni duplication en routant les paquets au mieux de ses possibilités (principe du *best effort*).

C'est également le principe de la poste qui s'efforce d'acheminer le courrier dans les meilleurs délais. Mais, pour s'assurer qu'un colis arrivera à destination en temps et en heure, il faut payer un service supplémentaire. De même, pour les réseaux, il faut mettre en œuvre des moyens supplémentaires afin de garantir ce niveau de service.

Alors que les flux de données classiques (Telnet, FTP, etc.) se contentent du service de base (celui du meilleur effort), le transport des flux audio et vidéo nécessite plus que cela. Il faut en effet garantir :

- qu'une application disposera du minimum de débit réseau nécessaire à son bon fonctionnement ;
- un temps de réponse ;
- que le temps de réponse variera peu dans le temps.

Afin de gérer la qualité de service nécessaire au traitement des flux multimédias, deux modèles sont actuellement proposés :

- La **différenciation de service** (modèle appelé **DiffServ**) repose sur l'affectation de priorités et de classes de service dont les valeurs sont transportées dans les paquets IP. Le flux est formaté (classé) à l'entrée du réseau, puis la qualité de service est appliquée de la même manière dans tous les routeurs en fonction de la valeur indiquée dans les paquets.
- L'**intégration de service** (modèle appelé **IntServ**) consiste à réserver les ressources tout le long du chemin qu'emprunteront les paquets, puis à appliquer à tout le flot de paquets qui suivent la qualité de service demandée lors de la réservation.

La qualité de service selon DiffServ

Le modèle DiffServ est une norme en cours de spécification qui ne précise pas encore les classes de service. C'est pour cela que les implémentations actuelles ne traitent que de la priorité (champ *IP Precedence*).

Le champ TOS

Dès l'origine, les concepteurs des protocoles TCP/IP ont pensé à intégrer la notion de qualité de service dans un champ du paquet IP, appelé TOS (*Type of Service*). Ce champ était destiné à transporter des informations relatives à la priorité et à la classe de service, mais il n'a jamais été réellement utilisé, jusqu'à ce que la qualité de service soit d'actualité. Mais, entre temps, les besoins ont évolué, et différents groupes de travail au sein de l'IETF ont proposé de modifier la signification de ce champ.

LE CHAMP TOS (RFC 791 ET 1349)

Le champ **TOS** (*Type of Service*) consiste en un unique octet :

IP precedence	TOS	0
3 bits	4 bits	

Le champ **IP precedence** (bits 0 à 2) détermine une priorité allant de 0 à 7 :

- 0 = Routine
- 1 = Priority
- 2 = Immediate
- 3 = Flash
- 4 = Flash override
- 5 = Critical
- 6 = Internetwork control
- 7 = Network control

Les bits **TOS** (la RFC utilise malencontreusement le même terme TOS pour l'octet entier ainsi que pour ces 4 bits) indiquent la classe de service souhaitée :

Bit	Description	Signification des valeurs
3	minimize Delay	(0 = Normal delay ; 1 = Low Delay)
4	maximize Throughput	(0 = Normal throughput ; 1 = High throughput)
5	maximize Reliability	(0 = Normal reliability ; 1 = Reliability)
6	minimize monetary cost	(ajouté par la RFC 1349)

Ces 4 bits positionnés à 0 indiquent tout simplement le service normal (le *best effort*).

Le dernier bit est réservé à un usage futur.

Les algorithmes WFQ et WRED prennent en compte le champ *Precedence* (d'où le “W” pour *Weighted*). L'activation de ces files d'attente est donc particulièrement judicieuse au sein du réseau, car elles permettent de réguler le flux en fonction des priorités fixées à l'entrée de celui-ci.

L'exemple suivant, qui montre quelques valeurs par défaut, permet d'expliquer comment WRED interagit avec les priorités :

```
random-detect precedence 0 109 218 10
random-detect precedence 1 122 218 10
...
random-detect precedence 7 194 218 10
```


La première valeur (de 0 à 7) correspond à la priorité indiquée dans le champ *Precedence*. Les trois paramètres suivants indiquent quant à eux :

- le **seuil minimal**, en nombre de paquets dans la file d'attente, à partir duquel les paquets commenceront à être rejetés ;
- le **seuil maximal**, en nombre de paquets dans la file d'attente, à partir duquel les paquets seront rejetés au rythme indiqué dans le paramètre suivant ;
- la **fraction** de paquets rejetés lorsque le seuil maximal est atteint, par défaut un sur dix.

Afin de favoriser les paquets prioritaires, le seuil minimal est d'autant plus élevé que la priorité est grande. L'algorithme WRED calcule la taille moyenne de la file d'attente et la compare aux seuils.

Le seuil maximal est déterminé automatiquement en fonction du débit de l'interface et de la mémoire disponible, tandis que le seuil minimal correspond à une fraction du seuil maximal dépendant de la priorité (par exemple, 1/2 pour 0, 10/18 pour 1, 11/18 pour 2, etc.). Les commandes "*sh queueing*" et "*sh int random-detect*" permettent de visualiser les paramètres actifs.

Configuration des routeurs

Les routeurs sont de nos jours destinés à opérer sur le réseau WAN et donc à être enfouis au sein du réseau. En vertu de cette conception, seule la fonction de *policing* du modèle Diff-Serv est implémentée sur nos routeurs, les autres fonctions opérant en entrée.

À la différence du *Traffic shaping* qui régule un flux erratique, le *policing* veille au respect du flux selon le profil de flux qui lui est indiqué. Il peut alors **rejeter** les paquets non conformes ou les **marquer** pour une régulation ultérieure sur d'autres équipements au sein du réseau :

```
int e0
```

```
rate-limit output 5000000 32000 45000 conform-action transmit exceed-  
action drop
```

Le mot clé "*output*" indique que la classification du flux est appliquée en sortie de l'interface :

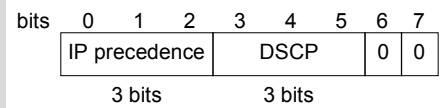
- le premier paramètre indique le débit moyen autorisé, ici 5 Mbit/s ;
- le deuxième précise la taille du burst, ici 32 Ko ;
- la troisième valeur est le burst en excès, ici 45 Ko.

Si le flux est conforme à ce profil (*conform-action*), le paquet est placé dans la file d'attente tel quel. Dans le cas contraire (*exceed-action*), il est rejeté.

LE POINT SUR DIFFSERV (RFC 2474 ET 2475)

Le modèle de différenciation de service repose sur le transport dans les paquets IP (champ TOS) de la **priorité** et de la **classe de service**. Il se propose en plus de définir la manière d'implémenter la qualité de service.

Le champ TOS est à l'occasion renommé **DS** (*Differentiated Service*) et structuré différemment, ce qui ne va pas sans poser de problèmes de compatibilité pour les routeurs traitant ce champ sur l'ancien mode (discussion objet de la RFC 2873).



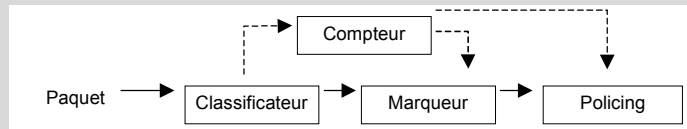
DiffServ conserve les 3 bits du champ Precedence et modifie les 3 suivants de manière à définir 64 **codepoints** répartis en trois pools :

- xx0 32 codepoints pour les actions standards ;
- x11 16 codepoints pour un usage expérimental et une utilisation locale ;
- x01 16 codepoints pour un usage expérimental, une utilisation locale ou une extension du premier pool.

Le codepoint xxx000 (issu du pool 1) permet d'assurer la compatibilité ascendante avec le champ Precedence. Les deux derniers bits 6 et 7 ne sont actuellement pas utilisés.

Le **classificateur** affecte une priorité et une classe de service au paquet en fonction de règles paramétrables (adresses IP source et destination, port TCP/UDP, etc.).

Le **conditionneur** applique la qualité de service en fonction du champ DS. Il analyse le flux de données (fonction de comptage — *meter*) et le compare à la QoS demandée. Il peut alors redéfinir la classe de service (fonction de marquage — *marker*), réguler le trafic (fonction de régulation — *shaper*) ou encore rejeter les paquets (fonction de rejet — *dropper*) pour l'adapter à la classe de service demandée.



À l'heure actuelle, deux classes de service sont définies :

- *Assured Forwarding* (RFC 2597), pour les flux nécessitant une bande passante limitée, le trafic en excès pouvant être rejeté progressivement selon un mécanisme de priorité à 12 niveaux (4 classes × 3 priorités de rejet).
- *Expedited Forwarding* (RFC 2598), également appelé "Premium service", pour les flux requérant une bande passante garantie avec des faibles taux de perte, de gigue et de latence.

Ces classes de service réservent des codepoints et décrivent la manière d'implémenter la gestion de la QoS au sein du réseau (*per-hop behavior*).

La fonction *policing* peut également marquer le paquet, c'est-à-dire lui affecter une priorité en fonction du respect ou non du profil de flux. On peut augmenter la granularité de cette dernière en précisant le protocole (et éventuellement les réseaux, voire les adresses source et de destination) sur lequel va s'appliquer le marquage :

```
rate-limit output access-group 101 5000000 24000 32000 conform-action
set-prec-transmit 1 exceed-action set-prec-transmit 0
```

```
access-list 101 permit tcp any any eq www
```

Nous avons décidé ici de limiter le flux des navigateurs web (protocole HTTP), à 5 Mbit/s. La priorité est fixée à 1 (prioritaire) s'il se conforme à la QoS, et à 0 (*best effort*) dans le cas contraire.

Sur notre backbone WAN à 2 Mbit/s, connecté en 100bT au LAN de notre site central, nous avons décidé de limiter le débit à 1 Mbit/s pour les paquets de priorité 1. Pour ceux de priorité 0, il est limité à 500 Kbit/s. Le burst autorisé est de 500 Ko, et de 500 Ko en excès.

interface 1/0

```
rate-limit input access-group rate-limit 100 1000000 500000 500000 conform-action transmit exceed-action drop
```

```
rate-limit input access-group rate-limit 101 500000 250000 250000 conform-action transmit exceed-action drop
```

```
access-list rate-limit 100 1
access-list rate-limit 101 0
```

Priorité (0 à 7) écrite dans le champ *IP Precedence*.

La commande “*access-list rate-limit*” est ici utilisée pour affecter une priorité basse si le volume du flux est important, et haute dans le cas contraire. Les différents équipements au sein du réseau se chargeront d'appliquer la QoS, par exemple, au niveau de leurs files d'attente WRED ou WFQ.

Configuration des commutateurs de niveau 2

En fait, il est bien plus judicieux de marquer les paquets à la source, c'est-à-dire sur les équipements directement en contact avec les PC. Il s'agit bien sûr des commutateurs Ethernet 10/100bT.

Les postes de travail sont, en effet, généralement connectés à des commutateurs de niveau 2, alors que l'on réserve la commutation de niveau 3 pour le réseau fédérateur en raison de son coût. Cependant, ils gèrent rarement les priorités. De ce fait, les commutateurs transforment les trames Ethernet en trames 802.1q pour gérer la qualité de service et les VLAN. Lorsqu'elles sont envoyées vers un port de sortie, ces trames sont reformatées en trames Ethernet 802.3.

Étant donné que toutes les trames ne sont pas obligatoirement de type 802.1q, par défaut le commutateur ne prend donc pas en compte le champ COS (*Class of Service*), et la priorité doit donc être fixée au niveau du port.

Par défaut, le port est en mode *untrusted* : la COS est fixée à la valeur par défaut du port, c'est-à-dire à 0. Il est possible de changer la valeur par défaut associée à chaque port :

```
set qos enable
set port qos 1/4 trust untrusted
set port qos 1/4 cos 6
```

Valeur par défaut du champ COS = 0

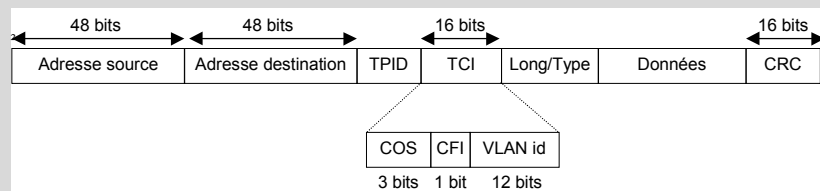
Fixe la valeur par défaut du champ COS

Toutes les trames arrivant par le port 1/4 auront un champ COS fixé d'office à 0, à moins qu'une autre valeur par défaut ne soit affectée au port.

Il n'est ainsi pas nécessaire de gérer la qualité de service au niveau du poste de travail et des serveurs. Chacun d'eux étant raccordé à un port du commutateur, celui-ci se charge de marquer les paquets. L'inconvénient est que ce marquage ne distingue pas le type de flux.

LA QUALITÉ DE SERVICE SUR ÉTHERNET (IEEE 802.1P)

La norme initiale 802.1q consiste à ajouter un champ à l'en-tête de la trame Ethernet initiale (802.3) à la fois pour gérer les VLAN et des **classes de service** (802.1p).



Cette trame est véhiculée uniquement entre les commutateurs. Ces derniers enlèvent le champ 802.1q lorsqu'ils transmettent la trame à un équipement terminal (PC, serveur, etc.).

Sept classes de service (**COS**) sont ainsi définies :

- 0 = Best effort
- 1 = Tâche de fond (batch)
- 2 = Réserve
- 3 = Excellent effort
- 4 = Applications à contrôle de charge (streaming audio/vidéo)
- 5 = Vidéo
- 6 = Voix
- 7 = Network control

Certains pilotes de carte réseau permettent de gérer la priorité directement au niveau du PC.

Si la priorité est déjà fixée par un équipement terminal ou un autre commutateur (donc générant une trame 802.1q), il est possible de prendre en compte le champ COS :

```
set qos enable
set port qos 1/3 trust trust-cos
```

Le champ COS des trames 802.1q arrivant par le port 1/3 sera ainsi pris en compte.

Configuration des commutateurs de niveau 3

Équipé d'une carte de commutation de niveau 3, le commutateur peut assurer toutes les fonctions DiffServ : classification, marquage et *policing*.

Par défaut, une politique de qualité de service est associée à un port, mais il est possible de l'associer à un VLAN :

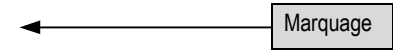
```
set qos enable
set port qos 1/1 port-based
set port qos 1/2 vlan-based
```

La configuration de nos commutateurs consiste à définir une **politique de qualité de service** contenant une règle de **marquage**, une règle de *policing* et une règle de **classification** (l'ordre de définition ne respecte pas celui du séquençement des opérations).

Définir une règle de marquage

Le commutateur associe aux ports une **politique par défaut** dont il est possible de modifier la règle de marquage :

```
set qos acl default-action ip dscp 0
```



Cette politique ne contient qu'une règle de marquage, qui consiste à marquer le champ DSCP (*Differentiated Service Code Point*) de tous les paquets avec la valeur indiquée (0 par défaut).

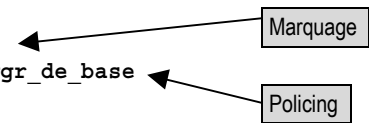
D'autres marquages sont possibles en remplaçant le mot clé *dscp* et sa valeur :

- *trust-dscp* : le champ DSCP n'est pas modifié et est accepté tel quel.
- *trust-ipprec* : le champ DSCP est modifié en prenant la valeur du champ *IP precedence*.
- *trust-cos* : le champ DSCP est modifié en fonction de la valeur contenue dans le champ COS de la trame Ethernet 802.1q.

Définir une règle de policing

Il est également possible d'associer plusieurs règles de *policing* à la politique par défaut :

```
set qos acl default-action ip dscp 0
microflow micro_de_base aggregate aggr_de_base
```



La création d'une règle de *policing* s'effectue comme suit :

```
set qos policier microflow micro_de_base rate 1000 burst 1000 drop
set qos policier aggregate aggr_de_base rate 1000 burst 1000 policed-dscp
```

Ces deux règles limitent le débit moyen à 1 Mbit/s et le dépassement (*burst*) à 1 Mo.

Le mot clé *microflow* précise que la règle s'applique à chaque flux considéré individuellement, tandis que le mot clé *aggregate* précise qu'elle s'applique à l'ensemble des flux.

Vient ensuite le type de traitement à appliquer au paquet lorsque le flux ne correspond pas au profil : le rejeter (*drop*) ou le marquer à nouveau (*policed-dscp*).

Dans ce dernier cas, la commande suivante permet de définir les correspondances, soit par plages de valeurs, soit individuellement :

```
set qos policed-dscp-map 63-62:62 61-60:60 ... 1-0:0
set qos policed-dscp-map 41:40
```

Les valeurs montrées ici sont celles par défaut qui sont conservées si elles ne sont pas explicitement redéfinies.

Définir une règle de classification

La politique par défaut qui vient d'être décrite ne contient pas de règle de classification puisqu'elle s'applique à tous les paquets IP. Par contre, nous pouvons créer une règle s'appliquant à un certain type de trafic :

```
set qos acl ip politique_ip_de_base trust-cos
microflow micro_de_base aggregate aggr_de_base
10.0.0.0 255.0.0.0 precedence routine
```

La politique ainsi définie s'applique aux paquets IP provenant du réseau 10.0.0.0 et ayant leur champ *precedence* à 0 (mot clé *routine*). Les paquets se conformant aux profils “*micro_de_base*” et “*aggr_de_base*” seront affectés de la priorité dérivée du champ COS de la trame 802.1q (mot clé *trust-cos*).

Les valeurs qui suivent le mot clé *precedence* sont celles décrites dans le champ TOS : *routine*, *priority*, *immediate*, *flash*, *flash-override*, *critical*, *internet* et *network*.

À la place de l'adresse IP et de son masque, on peut mettre le mot clé *any* ou le mot clé *host* suivi d'une adresse IP.

Nos commutateurs nous offrent la possibilité de décrire d'autres types de politiques s'appliquant aux trafics TCP, UDP, ICMP, IGMP, etc. Par exemple, il est ainsi possible d'affecter une règle à chaque application (qui se distingue par son port TCP ou UDP) :

```
set qos acl ip politique_web trust-ipprec
microflow micro_de_base
any tcp www
```

Associer une politique à un port

La dernière étape consiste à associer la politique que nous venons de créer à un port du commutateur ou à un VLAN :

```
set qos acl map politique_ip_de_base 1/1
set qos acl map politique_ip_de_base VLAN 1
```

Les ports ou le trafic associés au VLAN 1 se verront appliquer la politique IP de base que nous venons de définir.

Affecter des valeurs au champ DSCP

Par défaut, tous les ports fonctionnent sur le mode *untrusted*, ce qui signifie que les priorités (CoS et DSCP) des paquets y entrant sont ignorées.

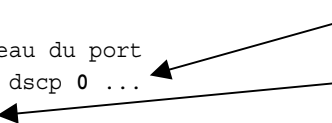
La carte de commutation de niveau 3 peut néanmoins prendre en compte les priorités des trames et paquets pour un port auquel est connecté une machine capable de générer des trames 802.1q et/ou des paquets IP DiffServ.

La manière de prendre en compte ces valeurs peut être définie soit au niveau de la règle de marquage de la politique qui est ensuite affectée au port comme nous l'avons déjà fait, soit être définie directement par port :

```
# Soit définie au niveau de la politique
set qos acl ip politique_ip_de_base trust-cos ...
set qos acl map politique_ip_de_base 1/1
```

```
# Soit définie directement au niveau du port
set qos acl ip politique_par_port dscp 0 ...
set port qos 1/4 trust untrusted
set port qos 1/3 trust trust-cos
set port qos 1/2 trust trust-ipprec
set port qos 1/1 trust trust-dscp
```

Valeur par défaut si
le port est marqué
untrusted



La signification des mots clés est la suivante :

- *untrusted* : le champ DSCP est fixé à 0 ou par la valeur indiquée dans la politique de qualité de service.
- *trust-cos* : le champ DSCP est fixé à la valeur du champ COS de la trame Ethernet.
- *trust-ipprec* : le champ DSCP est fixé à la valeur du champ *precedence* du paquet IP.
- *trust-dscp* : le champ DSCP du paquet est conservé tel quel.

Les correspondances COS (0 à 7) → valeur DSCP (0 à 63) et *IP precedence* (0 à 7) → DSCP par défaut conviennent, mais il est possible de les modifier comme suit :

```
set qos cos-dscp-map 0 8 16 24 32 40 48 56
set qos ipprec-dscp-map 0 8 16 24 32 40 48 56
```

Chaque valeur correspond à un DSCP compris entre 0 et 63. Les valeurs indiquées dans cet exemple sont celles par défaut.

Lorsque le paquet est finalement envoyé sur le port de sortie, le champ COS de la trame est dérivé de la valeur DSCP (conservée ou marquée à nouveau) :

```
set qos dscp-cos-map 0-7:0 8-15:1 ... 48-55:6 56-63:7
```

Là encore, les valeurs indiquées sont celles par défaut.

Configuration des postes de travail

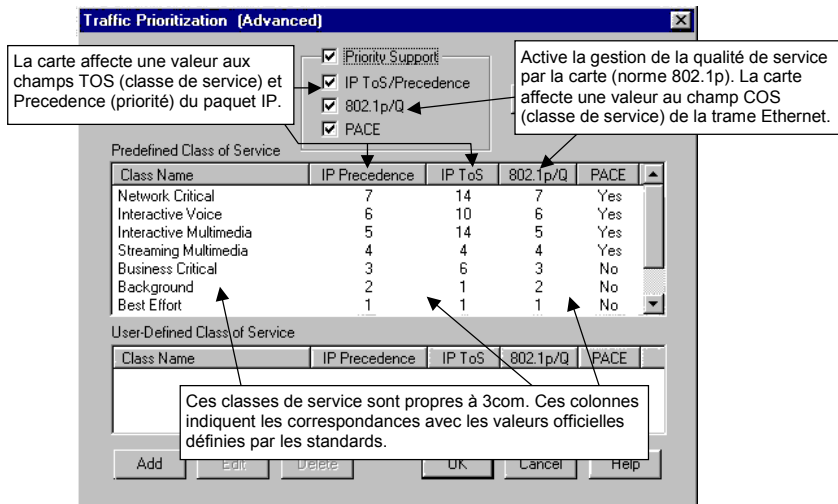
Nous venons de le voir, il est possible de gérer la qualité de service au niveau du WAN et du LAN. Cette gestion s'appuie sur des valeurs positionnées par les routeurs et les commutateurs de niveau 3 dans les paquets IP (champs TOS et Precedence), ainsi que par les commutateurs de niveau 2 dans les trames Ethernet (champ COS).

La question se pose alors de savoir où positionner ces paramètres. On peut le faire :

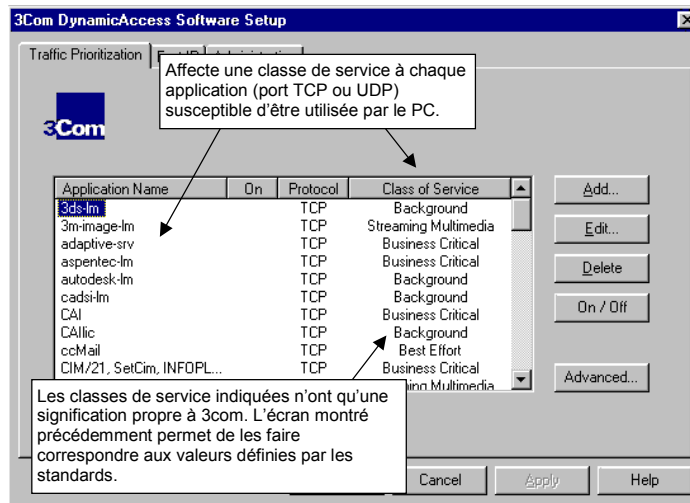
- au niveau des commutateurs de niveau 2 afin de gérer la qualité de service sur le LAN uniquement ;
- au niveau des commutateurs de niveau 2 et 3 afin de gérer la qualité de service sur le LAN et le WAN ;
- au niveau des routeurs pour gérer la qualité de service sur le WAN.

Cette dernière solution est préférable si les sites ne sont pas tous équipés de cartes de commutation de niveau 3. Par ailleurs, la bande passante est généralement suffisante sur le LAN pour permettre de se passer d'une gestion élaborée de la qualité de service.

Une dernière solution est proposée par les constructeurs : l'affectation des priorités au niveau des cartes Ethernet des postes de travail et des serveurs, c'est-à-dire à la source d'émission des trames et paquets.



Il faut donc affecter une classe de service à chaque application que sont susceptibles d'utiliser les postes de travail et les serveurs. Et il faut, bien évidemment, que toutes les cartes soient configurées de manière identique.



Cette solution est cependant difficile à mettre en œuvre pour plusieurs raisons :

- Tous les PC et serveurs doivent être équipés de cartes supportant cette fonctionnalité.
- Toutes les cartes doivent, de préférence, provenir du même constructeur, afin de simplifier l'exploitation et de garantir l'homogénéité des paramètres.

- Toutes les configurations doivent être identiques.
- Le paramétrage du driver doit pouvoir être réalisé à distance.
- Le nombre de nœuds à configurer est considérablement plus important que le nombre de commutateurs sur lesquels ce même paramétrage peut être défini.

Toutes ces contraintes rendent donc difficile la gestion de la qualité de service au niveau des postes de travail.

La qualité de service selon IntServ

À la différence de DiffServ qui repose sur le transport de la qualité de service dans les paquets IP, le modèle IntServ distribue ces informations (débit, temps de réponse, etc.) à tous les routeurs concernés par le flux, afin d'assurer la même qualité de service de bout en bout, c'est-à-dire entre les applications émettrices et destinataires. Cette cohérence nécessite deux fonctions distinctes :

- la définition des **spécificateurs** (paramètres) qui permettent de caractériser la qualité de service ;
- le protocole de **signalisation** permettant aux routeurs d'échanger les spécificateurs ainsi que des informations relatives à l'état du réseau.

La première fonction fait l'objet des RFC 2211 à 2216 qui définissent les éléments d'un réseau à intégration de service. La seconde est couverte par les RFC 2205 à 2209 qui définissent le protocole RSVP (*Resource reSerVation Protocol*).

La RFC 2210 établit la liaison entre les deux fonctions : elle traite du fonctionnement de RSVP au sein d'un réseau à intégration de service.

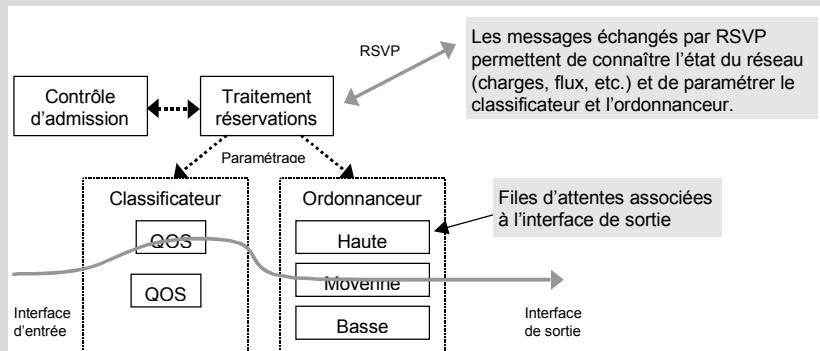
RFC	Sujet traité
2211	Classe de service par contrôle de charge
2212	Classe de service par garantie de service
2213	MIB de l'intégration de service
2214	MIB de l'intégration de service — Extension à la garantie de service
2215	Définition des paramètres permettant de calculer une QoS
2216	Définitions de la QoS pour les éléments du réseau (définit un cadre général sur lequel s'appuient les RFC 2211 et 2212)
2210	Utilisation de RSVP dans un réseau à intégration de service
2205	Spécifications fonctionnelles de RSVP
2206	MIB de RSVP
2207	Extensions pour les flux IPSEC
2208	Préconisations pour la mise en œuvre de RSVP
2009	Règles de traitement des messages

LE POINT SUR INTSERV (RFC 1633)

Il existe différentes façons d'assurer une qualité de service : affectation des priorités, choix de l'algorithme pour gérer les files d'attente, etc. Mais, si chaque routeur gère un niveau de priorité qui lui est propre ou gère des files d'attente selon des paramètres différents, la qualité de service demandée par une application sera interprétée différemment par les routeurs et ne pourra donc pas être assurée de bout en bout.

Pour que tout le monde parle le même langage, la RFC 1633 définit un cadre fonctionnel, appelé **intégration de service**. Quatre nouveaux composants doivent être intégrés aux routeurs :

- Le **contrôle d'admission** dont le rôle est de vérifier que la QoS demandée pour un nouveau flux ne remettra pas en cause les QoS accordées aux autres flux existants.
- Le **classificateur** de paquets dont le rôle est d'affecter une classe de service aux paquets entrant.
- L'**ordonnanceur** de paquets, qui gère les files d'attente pour les paquets sortant.
- Un mécanisme de **réservation** des ressources permettant de paramétrer le classificateur et l'ordonnanceur.



Une demande de réservation est traitée par le module de contrôle d'admission qui vérifie que l'application a le droit d'effectuer une telle requête et que le réseau peut offrir la QoS demandée sans remettre en cause celles déjà accordées.

L'implémentation logicielle de l'ordonnanceur dépend du routeur. La plupart d'entre eux utilisent néanmoins un algorithme standard de type WFQ (*Weighted Fair Queue*) ou WRED (*Random Early Detection*) pour la gestion des files d'attente.

De même, le classificateur peut utiliser différents éléments pour classer un paquet dans tel ou tel niveau de service : l'adresse IP et le port TCP/UDP sont des informations de base qui peuvent être utilisées.

Enfin, le mécanisme de réservation repose sur le protocole de signalisation **RSVP** (*Resource reSerVation Protocol*) qui permet de demander à tous les routeurs concernés de réserver des ressources.

La réservation des ressources

La signalisation RSVP est le protocole qui permet de négocier et de mettre en place une qualité de service dans des routeurs. Dans la pile TCP/IP, il est situé au même niveau que ICMP et IGMP (paquet IP de type 46, adresse unicast ou multicast identique à celle du flux). Il est cependant possible de l'encapsuler dans un paquet UDP (ports 1698 et 1699, adresse 224.0.0.14).

Les paquets RSVP sont routés saut par saut, c'est-à-dire de routeur à routeur : à chaque routeur traversé, les paquets remontent donc à la couche RSVP, puis sont réémis. Le routage des paquets est assuré par un protocole de routage quelconque, unicast ou multicast selon le type d'adresse.

Le protocole doit être activé sur chaque interface :

```
int s 0
ip rsvp bandwidth
```

Par défaut, 75% de la bande passante est géré par RSVP

Un émetteur de flux multimédia envoie périodiquement une annonce de chemin à un destinataire unicast ou à un groupe multicast (message *Path*). Chaque routeur, situé entre l'émetteur et le (ou les) destinataire, garde une trace de cette annonce avec les interfaces d'entrée et de sortie du paquet.

Chaque destinataire peut demander à réserver des ressources correspondant à une qualité de service telle que définie dans la RFC 2210 (message *Resv*).

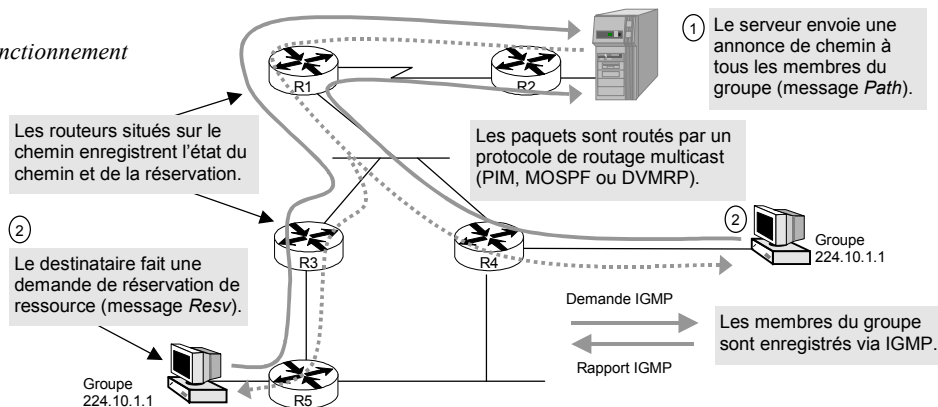
Le logiciel RSVP, qui reçoit une demande de réservation, procède aux traitements tels que spécifiés par la RFC 1633 :

- contrôle d'admission afin de garantir que la demande ne remettra pas en cause la QoS d'autres flux (l'application a-t-elle le droit de faire une telle demande ? Est-elle raisonnable ? etc.) ;
- affectation du flux à une classe de service permettant d'obtenir la QoS désirée.

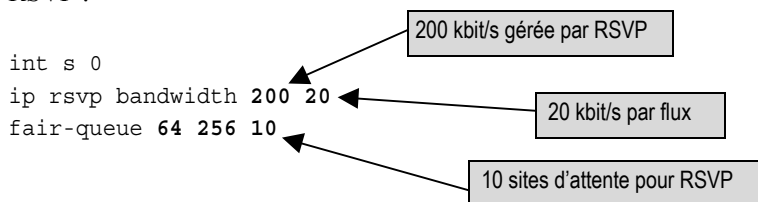
Si ces contrôles sont positifs, le message *Resv* est envoyé à l'émetteur en suivant la même route que le message *Path* émis. Chaque routeur situé entre le destinataire et l'émetteur du flux garde une trace de cette demande, et procède au même contrôle que décrit précédemment. Chaque routeur peut donc accepter ou refuser la demande (voir *figure 14-1* page 308).

RSVP ne définit que le moyen d'établir une qualité de service entre un émetteur et un destinataire en garantissant que les routeurs réserveront suffisamment de ressources (mémoire, CPU, files d'attente, bande passante sur les liaisons, etc.) pour assurer cette qualité de service.

Figure 14-1.
Principe de fonctionnement
de RSVP.



Sur nos routeurs, il est possible de définir quelle bande passante va pouvoir être gérée par RSVP :



L'algorithme WFQ a également été activé pour gérer la qualité de service demandée. Dix files d'attente sont ainsi dédiés à RSVP : c'est, en effet, le nombre de flux auquel nous nous attendons ($10 \times 20 \text{ Kbit/s} = 200 \text{ Kbit/s}$). Les valeurs par défaut, 64 et 256, ont été conservées pour les autres files d'attente.

LE POINT SUR RSVP (RFC 2205 À 2210)

L'émetteur d'un flux (vidéo, par exemple) envoie régulièrement aux destinataires un message **Path** leur permettant, ainsi qu'aux routeurs situés sur le chemin emprunté par le paquet RSVP, de déterminer la route empruntée par les paquets IP du flux. Les destinataires renvoient régulièrement à l'émetteur un message **Resv** qui est pris en compte par tous les routeurs situés sur le chemin qui vient d'être déterminé. Cette méthode permet de s'assurer que les ressources seront réservées sur le chemin emprunté par les paquets du flux multimédia. Au sein d'un réseau IP, le chemin retour peut, en effet, être différent du chemin aller.

Une demande de réservation consiste en un descripteur de flux composé d'un spécificateur de flux, **flowspec** (c.a.d. la QoS désirée qui permet de paramétrer l'ordonnancement de paquets) et d'un filtre de flux, **filterspec** (c.a.d. les caractéristiques du trafic qui permettent de paramétrer le classificateur de paquet).

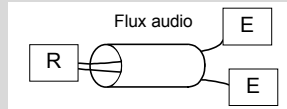
Un filtre consiste, au minimum, en une adresse IP et en un numéro de port TCP/UDP. Une forme plus évoluée peut inclure des données issues des couches applicatives ou le champ TOS du paquet IP.

...

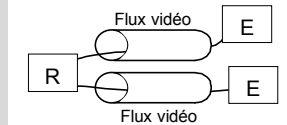
LE POINT SUR RSVP (SUITE)

Il existe plusieurs **styles de réservation** selon que la demande est partagée ou non par plusieurs flux et qu'elle concerne explicitement ou non les émetteurs :

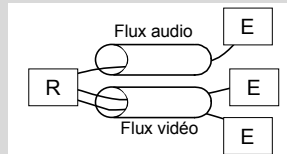
WF (Wildcard-Filter) : une seule réservation est faite pour les flux de tous les émetteurs (par exemple, pour une audioconférence où le nombre de personnes parlant en même temps est limité).



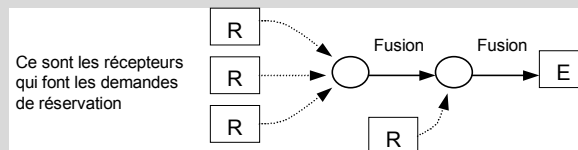
FF (Fixed-Filter) : une réservation est faite par flux (par exemple, une réservation par flux vidéo émis par chaque participant d'une visioconférence).



SE (Shared Explicit) : une seule réservation est faite pour les flux de quelques émetteurs (elle est adaptée, par exemple, à une visioconférence associée à une conférence audio).

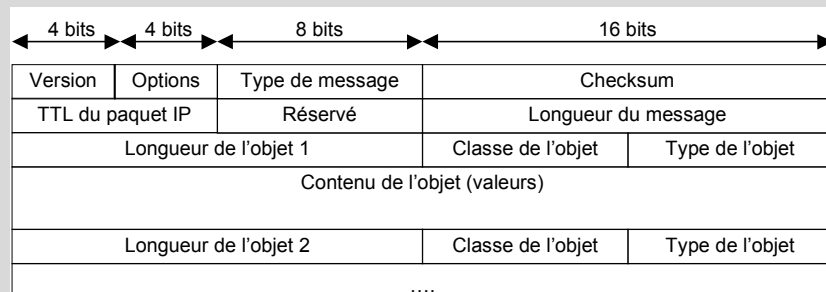


Un routeur peut recevoir des demandes de réservation en provenance de différents destinataires qui concernent un même flux émis par un même émetteur. Les QoS demandées seront alors fusionnées en une seule qui sera remontée, et ainsi de suite jusqu'à l'émetteur du flux. La réservation qui est remontée jusqu'à l'émetteur du flux dépend alors du style des réservations émises par les destinataires.



Un message RSVP comprend plusieurs objets, eux-mêmes structurés en plusieurs champs. Par exemple, un message *Resv* comprend les adresses IP de l'émetteur et du destinataire du message RSVP, la périodicité d'envoi du message et le style de réservation, suivi des descripteurs de flux.

Afin de tenir compte des changements de topologie (apparition et disparition des routeurs et des membres de groupe), les messages *Path* et *Resv* sont envoyés périodiquement.



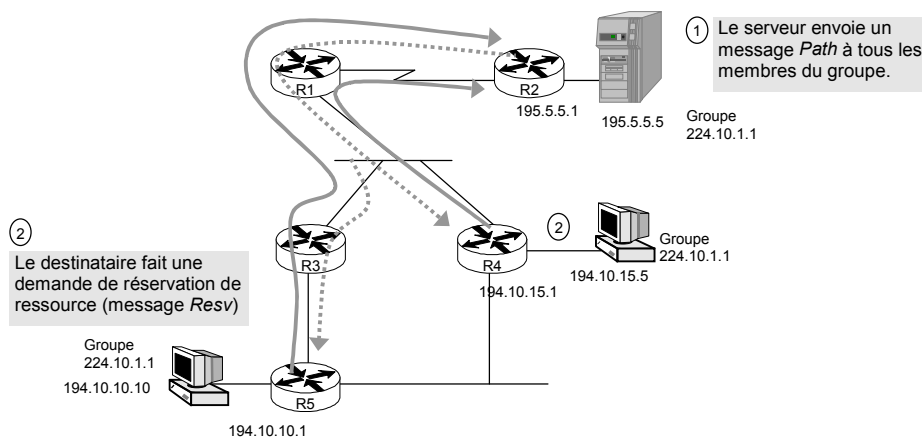
Ce qui vient d'être expliqué implique deux choses :

- que les piles TCP/IP des PC et des serveurs supportent le protocole RSVP ;
- et que les applications sachent décrire et demander une QoS.

Le premier point est résolu avec Windows 98, Windows NT 5.0 et Windows 2000. Le second est, quant à lui, résolu par les API RSVP disponibles dans Winsock 2.0. et sous la plupart des versions d'Unix. Ainsi, certaines applications, telles qu'Internet Explorer ou Netmeeting, prennent directement en charge ce protocole pour les flux audio et vidéo.

Mais peu d'applications prennent actuellement en charge RSVP et les spécificateurs de QoS. Heureusement, nos routeurs offrent une fonctionnalité intéressante qui consiste à simuler l'émission des messages *Path* et *Resv* par un serveur ou un PC.

Dans notre cas, qui sert d'exemple, un serveur doit diffuser une vidéoconférence à différents utilisateurs *via* le protocole Proshare d'Intel. L'architecture est la suivante :



La commande suivante permet de simuler la réception de messages *Path* par notre serveur vidéo :

```
#Routeur R2
ip rsvp sender 224.10.1.1 195.5.5.5 tcp 1652 4001 195.5.5.5 e0 20 5
```

Adresse du serveur supposé envoyer le message *Path*.

Adresse du destinataire du message (unicast ou multicast)

Le routeur agit comme si une application sur le serveur d'adresses IP 195.5.5.5 lui envoyait un message *Path* pour l'adresse de groupe multicast 224.10.1.1.

Les paramètres qui suivent sont :

- les ports destination (tcp 1652 = vidéoconférence Proshare) et source (4001) – les mots clés “udp” ou un port IP quelconque sont également acceptés ;
- l’adresse du saut précédent, donc celle du serveur ou du routeur le plus proche du serveur ;
- l’interface physique par laquelle est censé transiter le message émis par le serveur ;
- la bande passante, en Kbit/s, à réserver (ici 20 Kbit/s) ;
- Le dépassement autorisé (*burst*), en Ko (ici 5 Ko).

En retour, il est nécessaire de simuler l’envoi d’un message *Resv* par le destinataire du flux :

```
#Routeur R5
ip rsvp reservation 224.10.1.1 195.5.5.5 tcp 1652 4001 194.10.10.10 e0 SE
rate 20 5
```

Le routeur agit comme si une application sur le PC envoyait un message *Resv* pour l’adresse de groupe multicast 224.10.1.1 et vers le serveur 195.5.5.5.

Les paramètres qui suivent sont :

- les ports destination (1652 = vidéoconférence Proshare) et source – le mot clé “udp” ou un port IP quelconque sont également acceptés ;
- l’adresse du saut précédent, donc celle du PC ou du routeur le plus proche du PC ;
- l’interface physique par laquelle est censé transiter le message émis par le serveur ;
- le style de réservation : FF (*Fixed-Filter*), WF (*Wildcard-Filter*) ou SE (*Shared-Explicit*) que nous avons retenu pour notre flux vidéo ;
- la classe de service demandée (voir paragraphe suivant) : “rate” (QoS garantie) ou “load” (contrôle de charge) ;
- la bande passante, en Kbit/s, à réserver (ici 20 Kbit/s) ;
- le dépassement autorisé (*burst*), en Ko (ici 5 Ko).

Nous devons également nous assurer que les flux Telnet disposeront toujours de la bande passante suffisante et des temps de réponse corrects :

```
#Routeur R2
Message Path Émetteur Récepteur Flux telnet Émetteur
ip rsvp sender 195.5.5.5 194.10.10.10 tcp 23 1025 195.5.5.1 e0
```

```
#Routeur R5
Message Resv Émetteur Récepteur Flux telnet Récepteur
ip rsvp reservation 195.5.5.5 194.10.10.10 tcp 23 1025 194.10.10.1 e0 FF
load 32 1
```


Toutes les interfaces de sortie (par rapport au flux du serveur vers le PC) de tous les routeurs situés sur le chemin emprunté par le flux vidéo doivent pouvoir gérer la qualité de service demandée, ce qui est réalisé à *travers* la file d'attente WFQ. Si l'on ne considère que notre flux vidéo, cela donne :

```
# Routeur
int e1          (int e0 pour R2, in s0 pour R1, int e2 pour R4)
ip rsvp bandwidth 200 20
fair-queue 64 256 10
```

200 kbits/s gérés par RSVP
et 20 kbits/s affrétés par flux

Pour les routeurs situés au sein de notre réseau qui utiliseraient l'algorithme WRED, il est possible de modifier les seuils de rejet appliqués par défaut à RSVP :

```
random-detect precedence rsvp 205 218 10
```

Les paramètres montrés sont ceux par défaut ; leur signification est identique à celle donnée en page 296.

La description de la qualité de service

RSVP est un protocole qui permet aux applications et aux routeurs de mettre en place une qualité de service en échangeant des informations. Mais, comment une application peut-elle décrire la qualité de service dont elle a besoin ? Quels paramètres doit-elle fournir aux routeurs ? Comment les routeurs configurent-ils leur classificateur et leur ordonnanceur ?

La réponse à ces questions passe par le respect d'un langage commun, qui est décrit dans les RFC 2210 à 2216. Ce langage commun repose sur trois types de paramètres qui permettent de caractériser la qualité de service :

- la classe de service demandée ;
- les paramètres décrivant la classe de service ;
- les paramètres décrivant les caractéristiques du flux de données pour lequel cette demande de qualité de service est faite.

Ces paramètres sont regroupés dans des objets transportés dans les messages *Resv*.

Les classes de service

Le support d'une classe de service implique que les routeurs paramètreront leur classificateur et leur ordonnanceur de manière identique et que, en définitive, ils se comporteront de la même manière, afin d'assurer une qualité de service uniforme, celle qu'attend l'application.

Deux classes de service sont actuellement définies :

- Un service de **contrôle de charge** (RFC 2211). En prenant en charge ce type de service, le routeur garantit que la plupart des paquets seront routés sans erreur et que le délai de transit n'excédera pas un certain seuil pour la plupart des paquets routés.
- Un **service garanti** (RFC 2212). En prenant en charge ce type de service, le routeur garantit que le délai de transit des paquets ne dépassera pas un seuil fixé et que l'application disposera de la bande passante demandée.

Description des classes de service

La RFC 2215 définit, quant à elle, les paramètres utilisés pour calculer les classes de service :

- *non-is_hop* : information décrivant qu'un routeur ne gère aucune QoS (routeurs sans intégration de service).
- *number_of_is_hops* : nombre de sauts à intégration de service (c'est-à-dire de routeurs gérant une QoS).
- *available_path_bandwidth* : bande passante disponible le long d'une route (qui traverse plusieurs routeurs et plusieurs réseaux).
- *minimum_path_latency* : délai minimal de transit d'un paquet le long d'une route (dépend notamment de la bande passante et du délai de traitement des routeurs).
- *path_mtu* : taille maximale des paquets le long d'une route : c'est la plus grande MTU (*Maximum Transmission Unit* – 1 500 octets pour *Ethernet*) autorisée par les différents types de réseaux empruntés par un flot de données.

Caractéristiques des flux

La même RFC 2215 définit également les paramètres que l'application communique aux routeurs *via* RSVP (regroupés dans une structure appelée *token_bucket_spec*). Elle décrit les caractéristiques du flux pour lequel l'application fait une demande de classe de service. Il s'agit :

- du débit moyen du flux (*token rate*), en octets par secondes ;
- du dépassement de débit (*bucket size*) autorisé pendant une courte période, en d'autres termes le *burst*, exprimé en nombre d'octets ;
- du débit en pointe (*peak rate*), en octets par secondes ;
- de l'unité de contrôle minimale, en octets, qui correspond à la taille minimale des paquets. Le routeur contrôle en permanence le débit, et l'ajuste par rapport au *token-bucket (rate + size)* et à la taille des paquets. Les paquets dont la taille est inférieure à cette valeur seront traités, dans le calcul de la QoS, comme ayant cette valeur ;
- de la taille maximale des paquets en nombre d'octets.

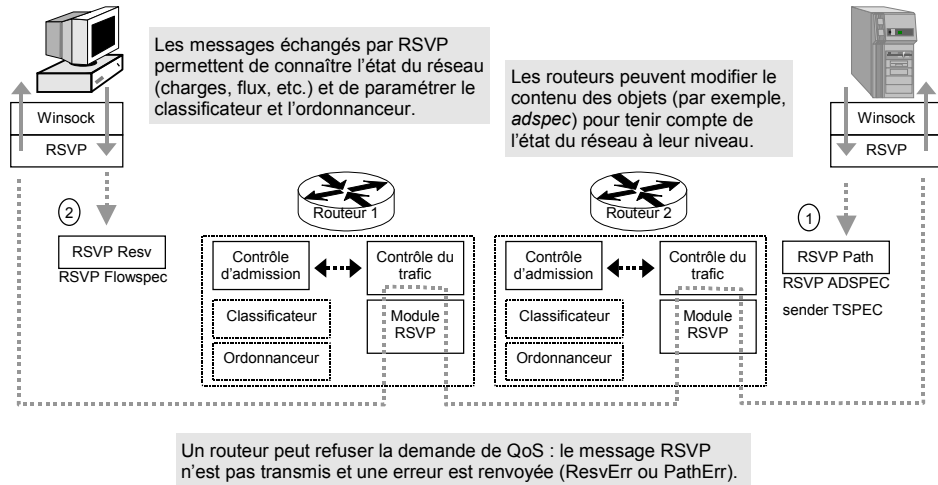
Objets RSVP

Ces paramètres sont regroupés dans des objets transportés dans les messages *Resv* de RSVP. Ces objets sont décrits dans la RFC 2210 :

- *flowspec* : est généré par le récepteur vers le ou les émetteurs d'un flux afin d'indiquer la QoS désirée. Les routeurs peuvent modifier les informations *flowspec* en cours de chemin, notamment lors de la fusion RSVP.
- *sender_spec* : indique les caractéristiques des flux que le ou les émetteurs vont générer. Les routeurs intermédiaires prennent en compte cette information sans la modifier.
- *rsvp_adspec* : est généré par le ou les émetteurs d'un flux, mais modifié par les routeurs intermédiaires. Cet objet contient des informations collectées dans le réseau, relatives aux délais de transit dans chaque nœud, à la bande passante mesurée sur les liaisons, à la classe de service supportée, etc. Chaque routeur modifie ou ajoute ses propres informations dans ce paquet.

Figure 14-2.

RSVP avec l'intégration de service (RFC 2210).



L'objet *flowspec* contient les paramètres suivants :

- la classe de service demandée : garantie ou contrôle de charge ;
- *tspec* (*traffic specification*) : correspond aux paramètres *token_bucket_tspec* ;
- *rspec* (*request specification*) : si la classe de service demandée est de type garantie. Cet objet décrit les caractéristiques de la classe de service, c'est-à-dire le débit (exprimé en paquets par seconde) et la gigue maximale autorisée (exprimée en microsecondes).

L'objet *sender_tspec* contient uniquement les paramètres *token_bucket_tspec*.

L'objet *rsvp_adspec* contient tout ou partie des paramètres de description de classe de service. Ils sont modifiés ou ajoutés par les routeurs pour décrire des caractéristiques locales et, en définitive, décrire les caractéristiques globales pour la route.

Définir une politique de qualité de service

L'application manuelle d'une politique de qualité de service sur l'ensemble des routeurs et des commutateurs peut se révéler fastidieuse et être source d'erreur.

Il est possible d'automatiser ce processus grâce au protocole COPS (*Common Open Policy Service*). La politique de qualité de service est alors définie à l'aide d'une interface graphique au niveau d'un serveur COPS, puis diffusée à partir de celui-ci aux routeurs du réseau. La création des règles et leur diffusion sont ainsi simplifiées.

Il faut avant tout indiquer aux routeurs de télécharger leur politique en utilisant le protocole COPS à partir d'un serveur :

Sur un routeur :

```
ip rsvp policy cops servers 10.0.20.2
```

Indique l'adresse IP du serveur COPS

Sur un commutateur :

```
set qos policy-source cops
```

```
set port qos 1/1 roles regulation
```

```
set cops server 10.0.20.2 diff-serv
```

Profil à télécharger pour ce port

Sur un autre commutateur :

```
set qos policy-source cops
```

```
set cops server 10.0.20.2 rsvp
```

```
set cops domain-name domaine_multimedia
```

QoS DiffServ ou IntServ

Nom du domaine défini au niveau du serveur COPS

Concernant le serveur lui-même et son interface graphique, il existe différents produits, tels que *QoS Policy Manager* (Cisco), *Orchestream* (Network Computing), *PolicyXpert* (HP) et *RealNet Rules* (Lucent).

15

La téléphonie et la vidéo sur IP

On entend souvent la conversation suivante entre un sceptique et un convaincu à propos de la téléphonie sur IP :

— Est-ce que ça marche ?

— Oui.

— Mais, est-ce que ça marche *bien* ? Est-ce que la qualité de service est bonne ?

La réponse est : oui, ça marche, et même très bien !

Maintenant que vous avez préparé votre réseau au multimédia, vous pouvez, en effet, mettre en place des solutions de téléphonie et de visioconférence.

Dans ce chapitre, vous apprendrez ainsi :

- à interconnecter des PABX *via* IP ;
- à raccorder un VoIP au réseau téléphonique classique ;
- le fonctionnement des protocoles multimédias.

Les protocoles H.323 de l'ITU-T sont ici expliqués dans le détail, car tous les produits du marché reposent sur cette norme. Mais, à l'avenir, cette dernière pourrait être remplacée par son concurrent SIP (*Session Initiation Protocol* — RFC 2543) promu par l'IETF. Les deux normes pourraient éventuellement cohabiter, SIP se posant en interface applicative, donc de plus haut niveau que H.323.

Présentation des protocoles multimédias

Historiquement le premier, le RTC (réseau téléphonique commuté) est le réseau qu'utilise toujours notre bon vieux téléphone. La version numérique de celui-ci, le RNIS (réseau numérique à intégration de service), est venue s'ajouter par la suite.

L'évolution des technologies a ensuite permis d'envisager d'autres supports pour transporter la voix : ATM, Frame Relay et, aujourd'hui, les réseaux IP (votre intranet et l'Internet). Toujours grâce à l'évolution des technologies, il est ensuite devenu possible d'y ajouter l'image, la vidéo et le partage de données.

Accompagnant le mouvement, les normes H.32x de l'ITU-T (*International Telecommunication Union — Telecommunication Standardization Sector*), sur lesquelles reposent le RTC et le RNIS, ont été adaptées à ces nouveaux supports de transmission et aux nouveaux besoins multimédias. Cette continuité dans la standardisation des protocoles téléphoniques permet ainsi d'assurer une cohabitation et une transition en douceur.

Réseau	Caractéristiques du réseau	Norme ITU-T
RNIS (réseau numérique à intégration de service)	Commutation de circuit Voix numérique	H.320
ATM (<i>Asynchronous Transfer Mode</i>)	Commutation de cellules Voix numérique	H.321
Réseaux WAN avec qualité de service	Frame Relay Voix numérique	H.322
Réseau VoIP (<i>Voice over Internet Protocols</i>)	Commutation de paquets IP Voix numérique	H.323
RTC (Réseau Téléphonique Commuté)	Commutation de circuit Voix analogique	H.324

Ces normes décrivent un cadre général de fonctionnement et préconisent l'utilisation d'autres normes de l'ITU-T et de l'IETF (*Internet Engineering Task Force*) en fonction du support de transmission (IP, RTC, RNIS, etc.). Elles décrivent les codages audio et vidéo, l'adaptation au support de transmission, la signalisation, etc.

	H.320	H.321	H.322	H.323	H.324
Réseau	RNIS	ATM	WAN QoS	VoIP	RTC
Codec vidéo	H.261 H.263	H.261 H.263	H.261 H.263	H.261 H.263	H.261 H.263
Codec audio	G.711 G.722 G.728	G.711 G.722 G.728	G.711 G.722 G.728	G.711 G.722 G.728 G.723 G.729	G.723
Adaptation	H.221	H.222	H.221	H.225	H.223
Signalisation	H.242 H.230	H.242	H.242 H.230	H.245 H.225	H.245
Conférences multipoints	H.231 H.243	H.231 H.243	H.231 H.243	H.323	
Données	T.120	T.120	T.120	T.120	
Couche de transport	I.400	I.363 AAL	I.400	TCP/IP	Fils

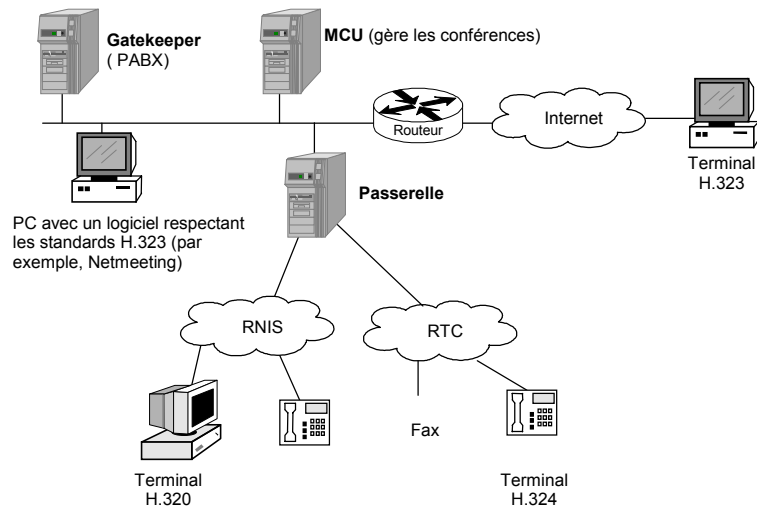
Ainsi, la norme **H.323** décrit un cadre général de fonctionnement pour un terminal multimédia fonctionnant sur IP. Elle implique l'utilisation d'autres normes décrites, cette fois-ci, par des RFC (*Request For Comments*) de l'IETF. Il s'agit notamment de **RTP** (*Real-time Transport Protocol*) qui décrit le format d'un paquet transportant des données (un échantillon sonore, une portion d'un écran vidéo, etc.).

Les composants d'un système H.323

La norme H.323 décrit le fonctionnement et l'interaction de quatre entités :

- un **terminal** qui supporte la voix et, optionnellement la vidéo et les données ;
- une **passerelle** qui permet l'interconnexion avec les autres réseaux H.32x tels que le RTC ;
- un **serveur de conférence**, appelé MCU (*Multipoint Control Units*) ;
- un PABX IP, appelé **gatekeeper**, qui offre le routage, la conversion d'adresses ainsi que la coordination de l'activité de toutes les entités H.323.

Figure 15-1.
Entités décrites
par la norme H.323.



La **passerelle** permet d'interconnecter le réseau téléphonique IP à d'autres réseaux tels que le RTC ou le RNIS. Elle assure la conversion des codecs audio et vidéo, de la signalisation et du support de transmission.

Le **MCU** contrôle l'entrée et la sortie des participants à la conférence, rediffuse le flux entre émetteurs et récepteurs en minimisant le trafic réseau, et assure l'éventuelle conversion de codec (par exemple, si un participant est équipé d'un écran de moins bonne qualité que les autres). Même si son utilisation n'est pas nécessaire lorsque les terminaux disposent des fonctions multicast et de négociation de paramètres, un MCU permet cependant de décharger le terminal de certaines fonctions.

Le **gatekeeper** gère des services pour les entités de sa **zone** de couverture, tels que :

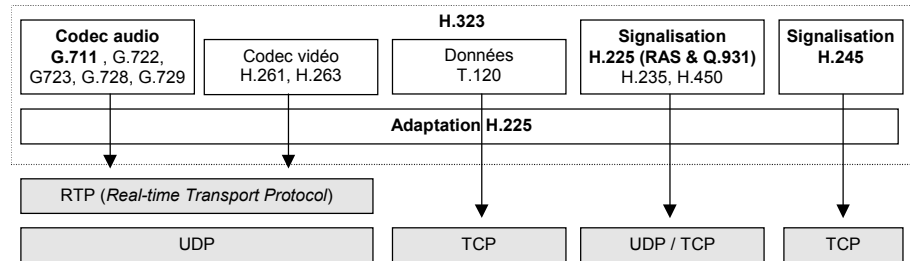
- la gestion des tables de correspondance entre les noms des terminaux (le nom de la personne), un numéro de téléphone E.164, une adresse e-mail et les adresses IP ;
- le contrôle d'admission : autorisation ou non de l'ouverture d'une communication ;
- la gestion de la bande passante sur le réseau, le nombre maximal de conférences, d'utilisateurs, etc. ;
- la localisation des passerelles et des MCU ;
- le routage des appels au sein de la zone et entre les zones.

Le **terminal**, enfin, se présente sous des formes diverses. Il peut prendre l'aspect d'un téléphone classique équipé d'une interface Ethernet dont la prise RJ45 est connectée à un commutateur de notre réseau local. Il peut également être un PC doté d'un logiciel de communication tel que Netmeeting de Microsoft. Le terminal intègre alors l'image, en plus de la voix, ce qui permet aux utilisateurs d'établir des visioconférences directement entre eux. Le PC devient ainsi un terminal multimédia traitant la voix, les données et l'image.

Le schéma suivant précise ce que la norme H.323 couvre (carrés blancs) et les éléments obligatoires (en gras). Les éléments grisés correspondent aux protocoles spécifiés par les RFC sur lesquels s'appuie la norme H.323 pour transporter les informations sur un réseau IP.

Figure 15-2.

Protocoles utilisés par la norme H.323.



La norme H.235 recouvre les fonctions de sécurité (authentification, intégrité des données, chiffrement, etc.). La norme H.450 décrit, quant à elle, les services complémentaires (identification de l'appelant, transfert d'appel, rappel sur occupation, etc.).

Selon ce schéma, un terminal H.323 doit donc comprendre :

- un codec audio (au moins G.711) ;
- optionnellement, un codec vidéo (au moins H.261) ;
- optionnellement, l'échange des données (à la norme T.120) ;
- les protocoles de signalisation RAS, Q.931 et H.245 ;
- une couche H.225 qui assure l'adaptation des protocoles de l'ITU à ceux de l'IETF ;
- les protocoles TCP, IP, UDP et RTP.

Par ailleurs, une entité H.323 comprend deux composants fonctionnels :

- un **MC** (*Multipoint Controller*) pour négocier, via H.245, les niveaux de service entre les terminaux et les ressources utilisées au sein d'une conférence (canaux audio et vidéo, adresses multicast, etc.) ;
- optionnellement, un **MP** (*Multipoint Processor*) pour traiter les flux multimédias (mixage, synchronisation de la parole et des images, diffusion, chiffrement, etc.).

En pratique, toutes les entités H.323 intègrent un MC pour participer à des conférences, tandis qu'un MCU intègre les deux composants pour gérer lesdites conférences.

L'établissement d'une communication

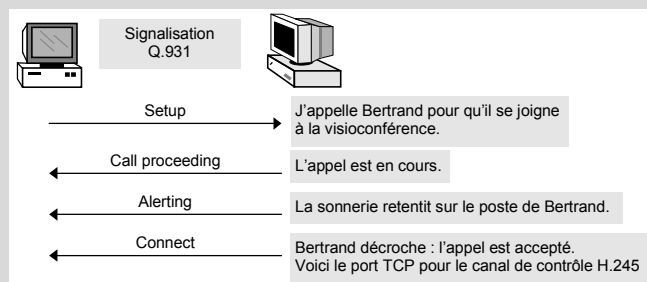
Pour communiquer entre elles, les entités H.323 ouvrent des canaux logiques (sessions TCP, UDP et RTP) soit dédiés à la signalisation, soit dédiés au transport des flux multimédias. On trouve, dans l'ordre :

- un canal de **signalisation RAS** (*Registration, Admission and Status*) qui permet à un terminal de s'enregistrer auprès du gatekeeper de sa zone ;
- un canal de **signalisation d'appel Q.931** (numérotation, sonnerie, etc.) qui permet à un terminal d'en appeler un autre ;
- un canal de **contrôle H.245** qui permet d'échanger les fonctionnalités supportées par les entités (codecs audio et vidéo, T.120) ainsi que d'ouvrir et de fermer les canaux audio, vidéo et données ;
- les canaux audio, vidéo et données.

L'utilisation des services d'un gatekeeper, et donc l'ouverture d'un canal RAS, est optionnelle. Les entités H.323 peuvent, en effet, communiquer directement entre elles.

LA SIGNALISATION D'APPEL Q.931

Une entité désirant en appeler une autre doit ouvrir un canal de signalisation d'appel Q.931 (dont l'utilisation au sein de H.323 est décrite par la norme H.225). Il s'agit d'une connexion sur le port TCP 1720 qui véhicule des messages Q.931.



Un message Q.931 comporte un en-tête, suivi d'un certain nombre d'éléments d'information obligatoires ou non. Par exemple, le message "Connect" contient, en plus de l'en-tête, les éléments d'information suivants : "Bearer capability", "Connect-UIIE", et optionnellement "Progress indicator", "Notification Indicator", "Facility", etc.

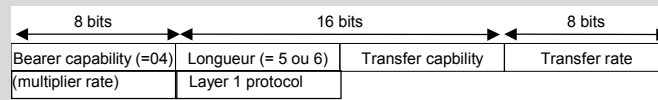
8 bits		16 bits		8 bits	
Discriminator (=08)		Référence d'appel (identifiant local)		Type de message	
Élément d'information		Champs de longueurs variables			
Élément d'information		Champs de longueurs variables			
etc.		etc.			

LA SIGNALISATION D'APPEL Q.931 (SUITE)

La signification des champs de l'en-tête Q.931 est la suivante :

- “ Type de message ” : 1 = Alerting, 2 = Call proceeding, 5 = Setup, 7 = Connect, etc.
- “ Éléments d'information ” : 4 = Bearer capability, 204 = numéro de téléphone de l'appelant, 224 = numéro de téléphone à appeler, 130 = temps de transit, etc.

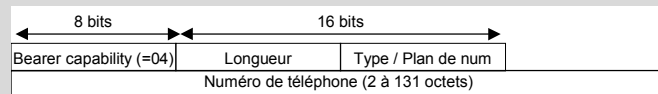
Par exemple, le format de l'élément d'information “ Bearer Capability ” est le suivant :



La signification des champs est la suivante :

- “ Transfer capability ” indique les fonctionnalités du terminal (audio, vidéo).
- “ Transfer rate ” indique le débit (mode paquet ou circuit, de 64 Kbit/s à 1 920 Kbit/s).
- “ Multiplier rate ” est uniquement présent si le champ précédent indique un débit “ multirate ”.
- “ Layer 1 protocol ” : G711 (A-law ou μ -law) ou H.225/H.245 (vidéophone H.323).

Voici un autre exemple montrant le format de l'élément d'information “ Called Number ” :



La signification des champs est la suivante :

- Type de numéro (3 bits) : 1 = international, 2 = national.
- Plan de numérotation (5 bits) : 1 = E.164, 8 = national, 9 = privé.
- Numéro de téléphone : par exemple 1#331836 ou 440553.

L'exemple suivant montre l'élément d'information utilisateur **Connect-UUIE** (*User-to-User Information Element*) pour le message “ Connect ” décrit selon la syntaxe ASN.1 (qui est utilisée dans toutes les normes de l'ITU-T) :

```
h323-message-body connect :{
  protocolIdentifier          { 0 0 8 2250 0 3 },
  h245Address                ipAddress : { ip '0A0C0006'h port 011026 },
  destinationInfo            { terminal { }, mc FALSE, undefinedNode FALSE },
  conferenceID               5A 1E 55 55 8A B0 CF 72 00 00 00 00 17 00 86 B4,
  callIdentifier              CallIdentifier,
  h245SecurityMode            H245Security (champ optionnel),
  tokens                     Sequence of ClearToken (champ optionnel),
  cryptoTokens               Sequence of CryptoH323Token (champ optionnel),
  fastStart                  Sequence of octet string (champ optionnel) }
```

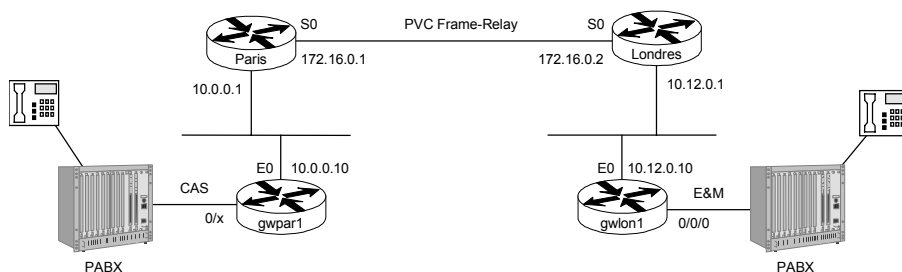
Interconnecter les PABX via IP

Lors de l'introduction d'une nouvelle technologie, la première contrainte est bien souvent de prendre en compte l'existant, en l'occurrence les PABX de l'entreprise.

Dans notre exemple, nous souhaitons profiter de notre réseau WAN pour acheminer les appels internes entre deux sites sur notre réseau privé (on Net), mais également pour passer des appels vers l'extérieur (off Net), de manière à bénéficier des meilleurs coûts (par exemple, Paris-Angleterre en sortant à Londres pour le coût d'un appel local ou national).

Figure 15-3.

Utilisation des passerelles pour interconnecter les PABX via IP.



Nous avons choisi des routeurs Cisco comme passerelles, mais cela aurait pu être des PC équipés de cartes de même nature que celles des routeurs.

La première tâche est de configurer les liaisons entre les passerelles et les PABX, ce qui est réalisé de la manière suivante :

```
#gwpar1
controler E1 0
framing crc4
linecode ami
mode cas
voice-group 1 timeslot 1-6 e&m-immediate
voice port 0/1
voice port 0/2
voice port 0/3
...
```

Configuration d'une liaison numérique CAS entre le PABX parisien et la passerelle gwpar1

6 canaux sont utilisés

```
#gwlon1
voice-port 0/0/0
signal immediate
operation 4-wire
type 2
voice-port 0/0/1
signal immediate
operation 4-wire
type 2
```

Configuration de deux liaisons analogiques E&M entre le PABX londonien et la passerelle gwlon1

Nous n'entrerons pas dans les détails, car ce domaine dépasse le cadre de cet ouvrage. Il faudrait, en effet, expliquer le fonctionnement des signalisations utilisées par les PABX classiques.

L'étape suivante consiste à configurer les liaisons entre les deux passerelles. Il s'agit cette fois de transporter la voix sur IP à l'aide des protocoles H.323.

Les utilisateurs des postes téléphoniques connectés aux PABX manipulent des numéros de téléphone au format E.164 (voir chapitre 10). Il faut donc associer des préfixes (c'est-à-dire la partie située le plus à gauche du numéro de téléphone) à des interfaces PABX d'un côté et à des interfaces IP de l'autre.

Passerelle	Préfixe	Correspondant
gwpar1 port 0/1-3	33	Numéro sur quatre chiffres
gwlon1 port 0/0/0	44	Numéro sur quatre chiffres

Chez Cisco, la connexion PABX est référencée par l'acronyme POTS (*Plain Old Telephone Service*), et la connexion IP est appelée VoIP (*Voice over IP*). Une connexion au sens large est appelée "dial-peer".

```
# gwpar1
dial-peer voice 33 pots
port 0/1
destination-pattern 33....

dial-peer voice 44 voip
destination pattern 44....
session-target ipv4:10.12.0.10
codec g729r8
```

Tous les numéros d'appels commençant par 33 sont envoyés sur cette interface.

Quatre points indiquent que quatre numéros doivent suivre le préfixe 33.

Tous les numéros d'appel commençant par 44 sont envoyés vers la passerelle de Londres en H.323.

Ce codec utilise 8 kbit/s de bande passante (cf. chapitre 11).

Sur la passerelle de Londres, nous devons disposer d'une configuration symétrique :

```
# gwlon1
dial-peer voice 44 pots
port 0/0/0
destination-pattern 44....

dial-peer voice 33 voip
destination pattern 33....
session-target ipv4:10.0.0.10
codec g729r8
```

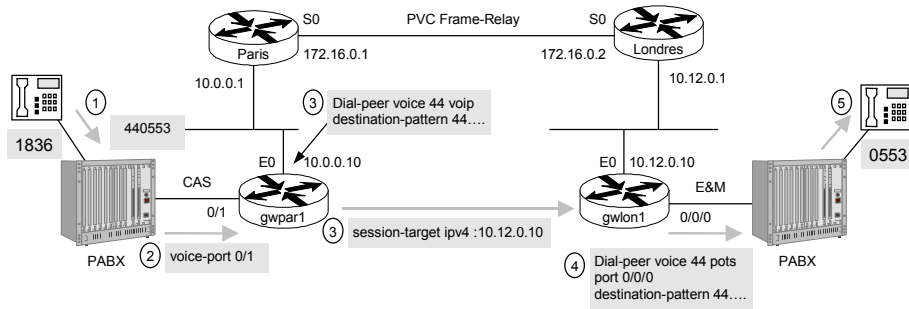
Numéro arbitraire ayant uniquement une signification pour la passerelle locale. Nous avons choisi une notation mnémotechnique.

Active les protocoles H.323

Adresse IP de la passerelle parisienne.

Figure 15-4.

Établissement d'une communication VoIP via des passerelles.



Le correspondant de Paris compose le “440553”. Le PABX route ce numéro vers la liaison CAS établie avec la passerelle parisienne. Celle-ci compare le numéro à ses “dial-peer” et constate que le préfixe “44” correspond à une connexion H.323. Elle envoie alors la demande de connexion sur IP (signalisation Q.931) à la passerelle londonienne qui compare à son tour le numéro présenté à ses “dial-peer”. Celle-ci constate alors que le préfixe “44” correspond à un port physiquement connecté au PABX et y envoie le numéro. Le PABX déclenche la sonnerie du téléphone du correspondant recherché.

Nous souhaitons également acheminer des appels en off Net sans que les utilisateurs français changent leurs habitudes de numérotation : “00” pour l’international, “44” pour l’Angleterre, suivis de dix chiffres. De même, les utilisateurs anglais numérotent comme suit : “00” pour l’international, “33” pour la France, suivis de neuf chiffres. Il suffit de configurer autant de “dial-peer” qu’il y a de préfixes :

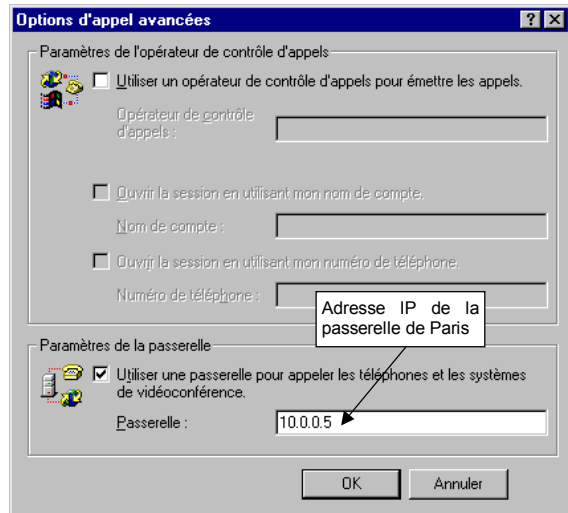
```
# gwpar1
dial-peer voice 0033 pots
port 0/4-6
destination-pattern 0033.....
dial-peer voice 0044 voip
destination pattern 0044.....
session-target ipv4:10.12.0.10
codec g729r8
```

Les neuf points correspondent aux neuf chiffres attendus.

Dix points = dix chiffres.

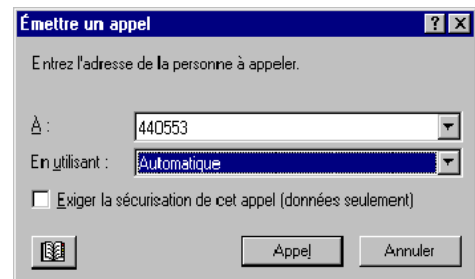
Un PC équipé de Netmeeting peut également profiter des services de nos passerelles pour appeler un correspondant situé sur le RTC classique. Il suffit d'indiquer son adresse IP dans le menu :

“Outil→Option→Appel Avancé”.



Il suffit ensuite de composer le numéro de téléphone de notre correspondant à Londres.

Nous utilisons ici la numérotation abrégée configurée dans nos passerelles.

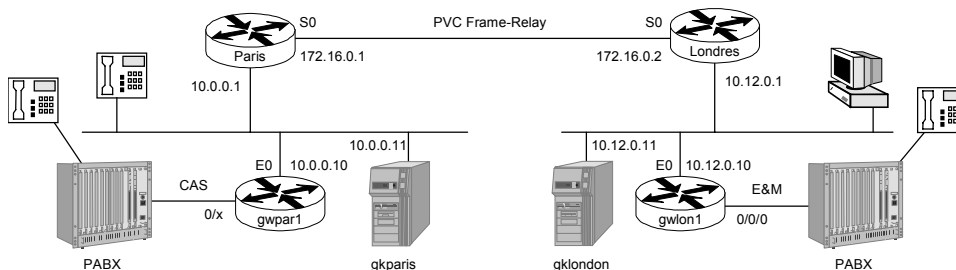


Mettre en place un gatekeeper

La complexité de notre installation augmente avec le nombre de sites à interconnecter. Nous devons de plus supporter des téléphones IP et des PC équipés de logiciels téléphoniques, tels que Netmeeting. Le réseau VoIP prenant de l'importance, nous devons assurer des services plus évolués, tels que le routage automatique, les statistiques ou encore la facturation. Nous devons donc assurer non seulement la reprise de l'existant mais également le développement de la voix tout IP.

La solution est de mettre en place un PABX IP, appelé **gatekeeper** dans la terminologie H.323. Nous choisissons donc d'en installer un par site, de marque Cisco (encore !), et reprenons les deux sites précédents comme exemple.

Figure 15-5.
Utilisation des gatekeepers.



Afin de simplifier les configurations et la vie des utilisateurs, nous choisissons également d'utiliser les services d'un DNS (*Domain Name System*) privé, dont nous avons déjà entrevu l'intérêt au chapitre 3 et dont la mise en œuvre est expliquée au chapitre 17. Le DNS permet de manipuler des noms à la place des adresses IP en assurant la résolution de l'un vers l'autre.

Dès lors, la configuration des gatekeepers est la suivante :

```
#gkparis
gatekeeper
zone local gkparis fr.intranet
zone prefix gklondon 44....
zone prefix gklondon 0044.....
gw-type-prefix 1# default-tech
```

Active H.323, la fonction de gatekeeper et la signalisation RAS.

Le gatekeeper gère les terminaux IP situés dans le domaine DNS indiqué.

Le gatekeeper route les appels préfixés 44 et 0044 vers Londres.

Le préfixe technique, indiqué par la commande "gw-type-prefix" — à ne pas confondre avec le préfixe de numérotation — permet de désigner explicitement les services d'une passerelle, par exemple une passerelle offrant le routage onNet, une autre le routage offNet, le

fax, etc. Par convention, le préfixe technique se termine par un dièse. Dans notre exemple, si aucun préfixe technique n'est indiqué dans le numéro appelé, l'appel sera routé vers la passerelle supportant le préfixe technique 1#.

Le gatekeeper de Londres est configuré de manière symétrique :

```
#gklondon
gatekeeper
zone local gklondon uk.intranet
zone prefix gkparis 33....
zone prefix gkparis 0033.....
gw-type-prefix 1# default-tech
```

Notre DNS privé est configuré avec un sous-domaine par pays, et nous installons un gatekeeper par zone.

Le préfixe technique par défaut est 1#.

Le routage des appels est désormais réalisé entre gatekeepers. Dans chacune des zones, les passerelles doivent par conséquent s'enregistrer auprès de leur gatekeeper de rattachement :

```
#gwparl
gateway
int e0
h323-gateway voip interface

h323-gateway voip h323-id gwparl@fr.intranet
h323-gateway voip id gkparis multicast
h323-gateway voip tech-prefix 1#
```

Active H.323 et la signalisation RAS.

Désigne explicitement l'interface Ethernet pour les fonctions H.323.

Nom de la passerelle dans la zone fr.intranet contrôlée par un gatekeeper

Recherche le gatekeeper gkparis à l'aide de la signalisation RAS dans sa version multicast.

À l'aide de ces commandes, la passerelle parisienne numéro 1 va rechercher le gatekeeper appelé "gkparis" et s'y enregistrer avec le préfixe technique 1# sous le nom "gwparl" dans la zone "fr.intranet".

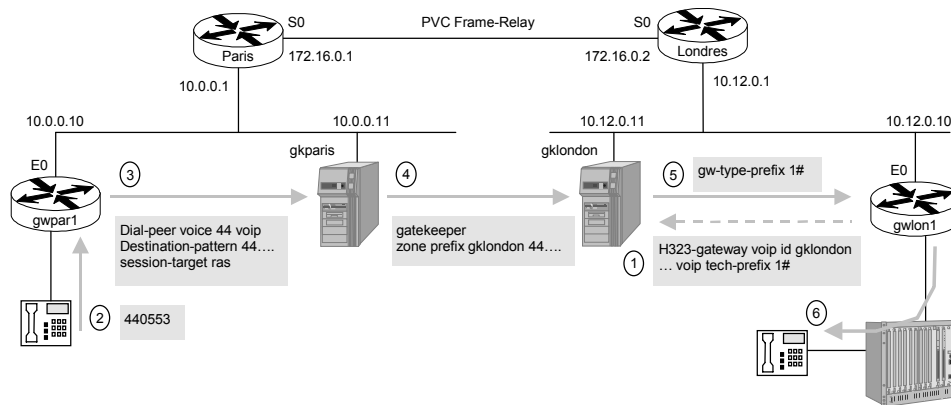
Toujours sur nos passerelles, la configuration des "dial-peer" change légèrement par rapport à une architecture sans gatekeeper :

```
dial-peer voice 44 voip
destination-pattern 44....
tech-prefix 1#
session-target ras
```

Optionnel : ajoute le préfixe 1# au numéro appelé, afin de désigner explicitement un groupe de passerelles.

Utilise les services du gatekeeper qui lui donnera l'adresse IP de la passerelle de Londres.

Lorsque ce “dial-peer” sera invoqué, il enverra l’appel au gatekeeper, qui le routera conformément à la configuration précédente.

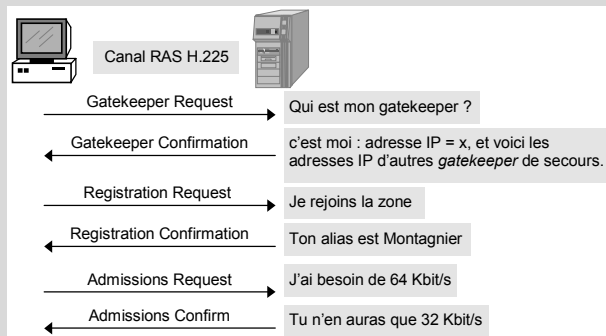


LE CANAL DE SIGNALISATION RAS (H.225)

La signalisation RAS (*Registration, Admission and Status*) permet à toutes les entités H.323 (un terminal, un MCU ou une passerelle) de communiquer avec un **gatekeeper** afin de bénéficier de ses services.

Les messages de **découverte** sont envoyés dans des paquets UDP, port 1718, soit à destination du gatekeeper s’il est connu (son adresse ayant été configurée manuellement), soit à destination du groupe multicast 224.0.1.41. Si plusieurs gatekeepers répondent, l’entité en choisit un.

L’étape suivante consiste à rejoindre la **zone** contrôlée par un gatekeeper en **s’enregistrant** auprès de ce dernier. Le canal RAS utilise maintenant des paquets UDP unicast, port 1719.



Désormais, l’entité H.323 s’adresse au gatekeeper pour appeler tous ses correspondants en les désignant par leur nom ou leur numéro de téléphone. Le gatekeeper se charge de convertir ce nom en une adresse IP, et de router l’appel :

- soit directement vers le correspondant si celui-ci se trouve dans la même zone ;
- soit vers un autre gatekeeper si le correspondant se trouve dans une autre zone ;
- soit vers une passerelle si le correspondant n’utilise pas un terminal IP (mais un téléphone, par exemple).

...

LE CANAL DE SIGNALISATION RAS (H.225 – SUITE)

Plusieurs types de messages peuvent être échangés : Découverte, Enregistrement/annulation, Admission (demande d'autorisation d'utiliser le réseau), Changement de bande passante (demandé par le terminal), Localisation (conversion d'adresses), Statut et Information sur les ressources disponibles (pour les passerelles). Ces types regroupent 18 messages différents.

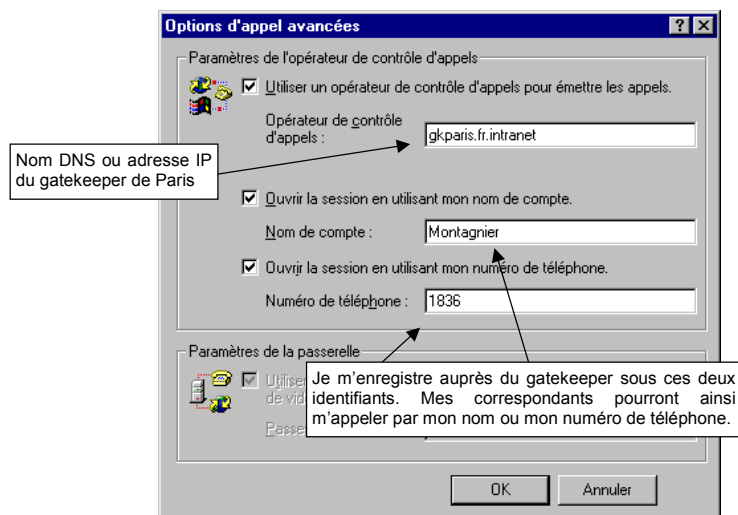
Voici, à titre d'exemple, le format, en syntaxe ASN.1, d'une demande d'enregistrement d'une passerelle auprès d'un gatekeeper :

```
RasMessage ::= registrationRequest : {
  requestSeqNum 037001,
  protocolIdentifier { 0 0 8 2250 0 3 },
  discoveryComplete TRUE,
  callSignalAddress { ipAddress : { ip 'A00C00A'h, port 01720 } },
  rasAddress { ipAddress : { ip A0 0C 00 0A, port 04520 } },
  terminalType { gateway { protocol { voice : { supportedPrefixes { { prefixe 164 : "1#" } } } }, mc FALSE, undefinedNode FALSE },
  terminalAlias { h323-ID : "gwlon1.fr.intranet" },
  gatekeeperIdentifier { "gklondon.fr.intranet" },
  endpointVendor { vendor { t35CountryCode 0181, t35Extension 00, manufacturerCode 018 } } }
```

ITU-T, recommandation
série H, 225.0, (0), version 3

Q.931

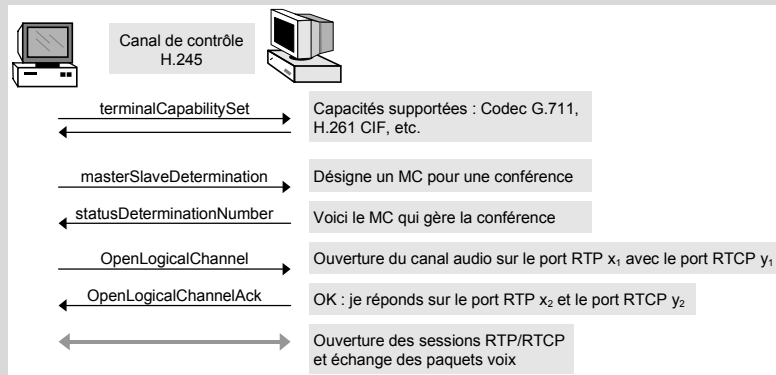
Un PC équipé de Netmeeting peut également utiliser les services d'un de nos gatekeepers. Il suffit d'indiquer son adresse IP, ou mieux, son nom DNS, dans le menu "Outil→Option→Appel Avancé".



LE CANAL DE CONTRÔLE H.245

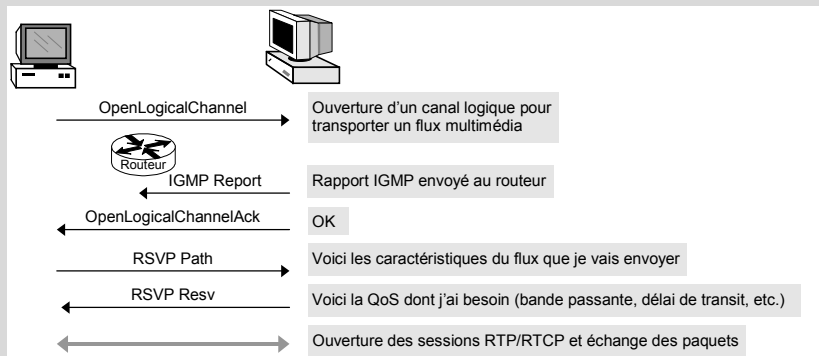
La dernière étape dans l'établissement d'une communication H.323 est l'ouverture du canal de contrôle H.245. Les messages échangés permettent de négocier les fonctions prises en charge par les terminaux (*Terminal Capability*) : choix des codecs audio et vidéo, de la résolution d'image, etc., puis d'affecter dynamiquement les ports UDP supportant les **canaux audio et vidéo**.

Si une conférence à trois ou plus est demandée, la procédure désigne le MC (celui d'un terminal, ou le MCU s'il existe) qui sera responsable de la gérer.

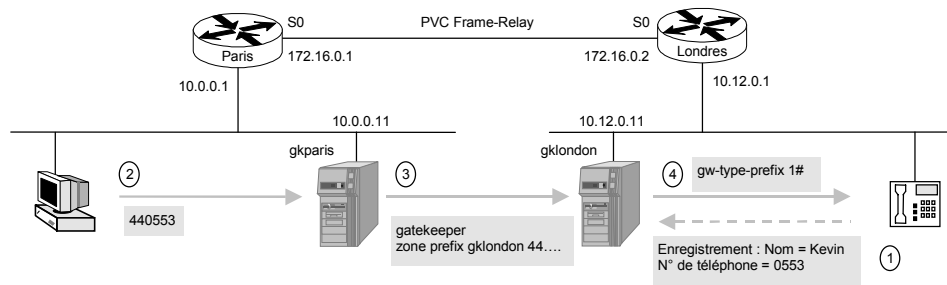


Le canal de contrôle H.245 véhicule ensuite les commandes permettant :

- d'ouvrir et de fermer les canaux audio et vidéo, c'est-à-dire les sessions **RTP** (*Real-time Transport Protocol*) ;
- de gérer les entrées et les sorties dans les conférences ;
- de gérer la **qualité de service**, grâce aux informations données par **RTCP** (*RTP Control Protocol*) et en faisant appel aux services de **RSVP** (*Resource Reservation Protocol*) pour réserver la bande passante nécessaire sur le réseau IP.

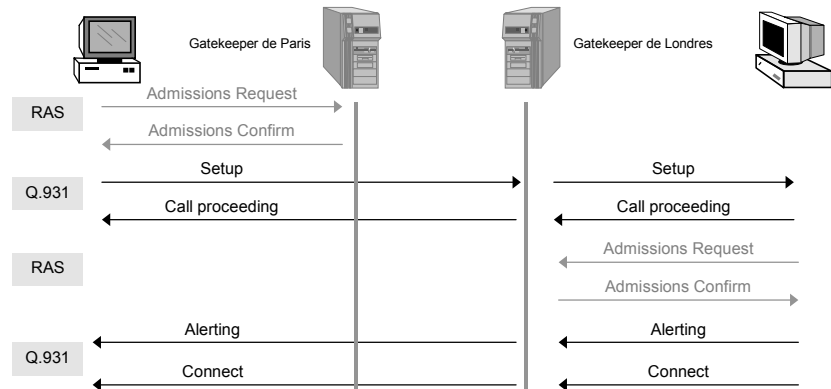


Si cela est nécessaire, un terminal peut demander au gatekeeper l'autorisation de changer de bande passante en lui envoyant un message "Bandwidth Change Request". S'il obtient l'accord, le canal de contrôle H.245 est fermé entre les deux terminaux, puis un nouveau est ouvert accompagné d'une nouvelle demande RSVP.



Sur chaque site, les terminaux s'enregistrent auprès de leur gatekeeper local (Paris et Londres dans notre cas). Lors d'un appel, le gatekeeper de Paris autorise l'ouverture directe du canal de signalisation d'appel entre le terminal et son correspondant, tandis que le gatekeeper de Londres impose que ce canal passe par lui.

Figure 15-6.
*Gestion des appel
via des gatekeepers.*



Le gatekeeper de Paris peut trouver les informations du correspondant concernant un terminal auprès d'un autre gatekeeper en utilisant le message "Resource Locator". Le gatekeeper de Londres lui renvoie alors son adresse ainsi que le port TCP sur lequel il répond. S'il avait décidé que la communication pouvait être directe, il aurait envoyé l'adresse IP du terminal.

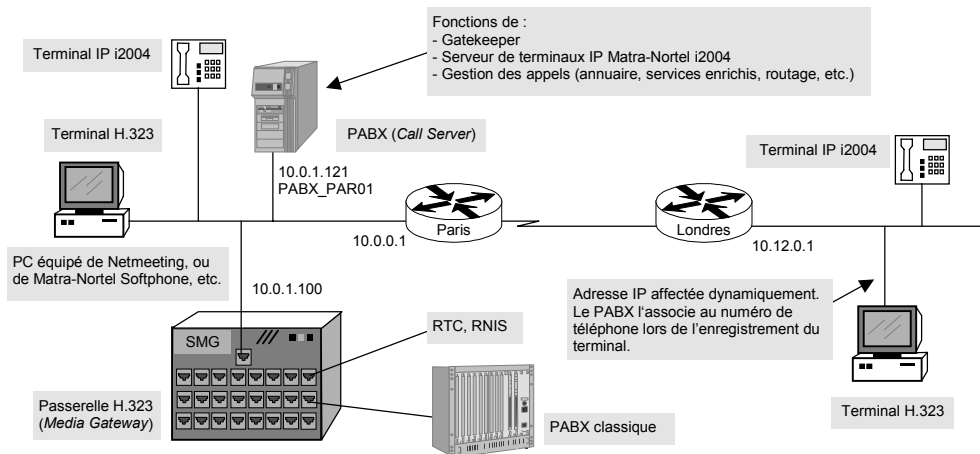
L'appel accepté, les deux entités peuvent ouvrir le canal de contrôle H.245 dans le but de négocier les paramètres du flux multimédia.

La voie vers le tout IP

Nous venons de mettre en place un système qui permet à nos terminaux IP (PC, téléphones IP) de communiquer entre eux (voix et image). Nous utilisons des gatekeepers pour assurer la conversion d'adresses et le routage des appels. Des passerelles sont également mises en place afin d'interconnecter notre réseau au monde téléphonique classique (PABX, RTC, RNIS, DECT, etc.).

Mais, par rapport à la téléphonie classique, il nous manque toute la gamme de services enrichis qu'un PABX peut offrir : plan de numérotation, messagerie vocale, transfert, filtrage, restriction et interception d'appels, etc.

L'aboutissement d'un réseau VoIP est donc la mise en place d'un PABX IP. Deux types d'offres sont proposées sur le marché : celles des constructeurs informatiques (Cisco, Lucent, etc.) et celles des constructeurs de PABX (Alcatel, Matra-Nortel, etc.). Les premiers ont l'avantage de bien connaître le monde IP ; les derniers offrent celui d'une plate-forme existante qu'il suffit de porter sur IP. C'est le cas de Matra-Nortel qui a porté sous Windows sa gamme 6500 fonctionnant à l'origine sous un Unix propriétaire. Le résultat est le produit Succession qui offre le même niveau de fonctionnalité qu'un PABX traditionnel.



Le PABX peut être situé n'importe où sur le réseau. Les sites importants en sont équipés, tandis que les petits sites n'en ont pas besoin, les terminaux dialoguant alors avec le PABX du site principal.

De son côté, la passerelle accueille les terminaux non IP : PABX classiques, téléphones analogiques et numériques Matra-Nortel, bornes DECT, etc. L'accès au monde extérieur (RTC, RNIS) est réalisé soit *via* le PABX traditionnel, soit directement par la passerelle. Cette dernière peut également accueillir les terminaux IP et peut, de ce fait, coupler le terminal VoIP à un poste DECT.

Le PABX accueille les terminaux H.323 natifs ainsi que ceux de Matra-Nortel *via* un protocole propriétaire. Ces derniers se comportent comme des terminaux passifs qui ne font qu'afficher les informations envoyées par le PABX (fonction équivalente à celle d'un serveur de terminaux, tel que MetaFrame).

Configurer le PABX et la passerelle VoIP

Pour configurer tous ces éléments, nous nous retrouvons dans un environnement familier d'une interface graphique sous Windows. Désormais, ces systèmes sont, en effet, considérés comme des services à part entière au même titre que les bases de données, les serveurs de fichiers, la messagerie, etc.

L'écran principal du gestionnaire fourni par Matra-Nortel présente les éléments de notre réseau VoIP : commutateurs du réseau local (ici un *BayStack 450*), passerelles (la *Media Gateway* de Nortel) et PABX (le *Call Server* de Nortel).

The screenshot displays the InfoCenter application window. The left sidebar shows a tree view of network resources, including 'Switches'. The main pane shows the 'Resources/Switches' configuration window, which lists three objects: '10.0.2.40' (Switch, BayStack 450), '10.0.1.100' (Switch, Succession - Media - Gateway), and 'PBX_PAR01' (Switch, Succession - Call - Server). A context menu is open for 'PBX_PAR01', showing options like 'Configuration', 'Performance', 'Weblinks', 'Open...', 'Cut', 'Copy', 'Net Aware Select', 'Device View', and 'Properties...'. The 'Configuration' option is selected, leading to a submenu with various OMSE (Operational Maintenance System) tasks, including 'Task scheduler', 'Subscriber management', 'Directory configuration', 'LCR Management', 'Graphical console', 'Telnet access', 'Provisioning', 'Explorer', 'Telephony features configuration', and 'Number range configuration'. Several callouts point to specific features: 'Paramètres de la ligne d'abonné' points to 'Task scheduler'; 'Gestion de l'annuaire' points to 'Directory configuration'; 'Gestion du routage des appels (Least Cost Routing)' points to 'LCR Management'; 'Gestion des services enrichis et des droits d'accès' points to 'Telephony features configuration'; and 'Gestion des alarmes sur le PABX' points to the 'Fault Summary' window at the bottom.

InfoCenter

File Edit View Options Admin Tools Window Help

InfoCenter

- Alarms
- Custom
- Network Seeds
 - All Seeds
 - User Seeds
- Resources
 - Bridges
 - ELANS
 - Hubs
 - Internet
 - Probes
 - Routers
 - Segments
 - Subnets
 - Switch Communities
 - Switches
 - VLANs
 - WANs

Resources/Switches

Objects: 3 Change Count: 3 Created: October 24, 2000 3:11 PM Last Modified: October 26, 2000 ...

Label	Type	SubType	Discovery St...
10.0.2.40	Switch	BayStack 450	
10.0.1.100	Switch	Succession - Media - Gateway	6
PBX_PAR01	Switch	Succession - Call - Server	

Context menu for PBX_PAR01:

- Fault
- Configuration
 - OMSE - Task scheduler
 - OMSE - Subscriber management
 - OMSE - Directory configuration
 - OMSE - LCR Management
 - OMSE - Graphical console
 - OMSE - Telnet access
 - OMSE - Provisioning
 - OMSE - Explorer
 - OMSE - Telephony features configuration
 - OMSE - Number range configuration
- Performance
- Weblinks
- Open...
- Cut
- Copy
- Net Aware Select
- Device View
- Properties...

Callouts:

- Paramètres de la ligne d'abonné
- Gestion de l'annuaire
- Gestion du routage des appels (Least Cost Routing)
- Gestion des services enrichis et des droits d'accès
- Gestion des alarmes sur le PABX

Le PABX est un élément du réseau comme un autre.

Fault Summary

File Edit View Options Help

Status	State	Severity	Type	Device	Create Date	T...
New	Clear	Low	Device Unreachable	10.0.1.100	Thu Oct 26 12:19:54 2000	2
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:36:17 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:34:13 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:34:07 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:34:03 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:32:25 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:31:45 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:27:31 2000	1
New	Aged	Low	High severity alarm on comp...	10.0.1.100	Tue Oct 24 16:01:43 2000	1

Faults Loaded 9Fault(s) NMS

Déclarer les terminaux téléphoniques

Ensuite, il convient de déclarer les terminaux rattachés à notre PABX : il peut s'agir de téléphones VoIP Matra-Nortel (i2004), de PC équipés de Netmeeting ou encore de téléphones classiques raccordés à une passerelle.

Nous entrons alors dans le monde des téléphonistes : il faut créer un **abonné** auquel on associe une entrée dans l'**annuaire** ainsi qu'une **fiche technique** qui décrit les caractéristiques du poste et les services auxquels l'utilisateur a droit.

Update subscription (N° 5090)

File Display Subscription Configuration

Subscriber number: 5090
 device: PABX_PAR01
 Slot:
 Last name: MONTAGNIER
 First name: JEAN-LUC

Subscriber type: IP subscriber - Etherset
 Model of set:
 Multiline: No

Technical record: Yes
 Complementary record: No
 Nb of complementary record: 0
 Nb of directory record: 1
 NM record: No

Annotations:
 - Numéro de téléphone (points to 5090)
 - Sur le PABX, le terminal IP peut être de type Matra-Nortel ou H.323.
 - Le terminal peut être rattaché à une passerelle ou au PABX (points to IP subscriber - Etherset)

Create technical record

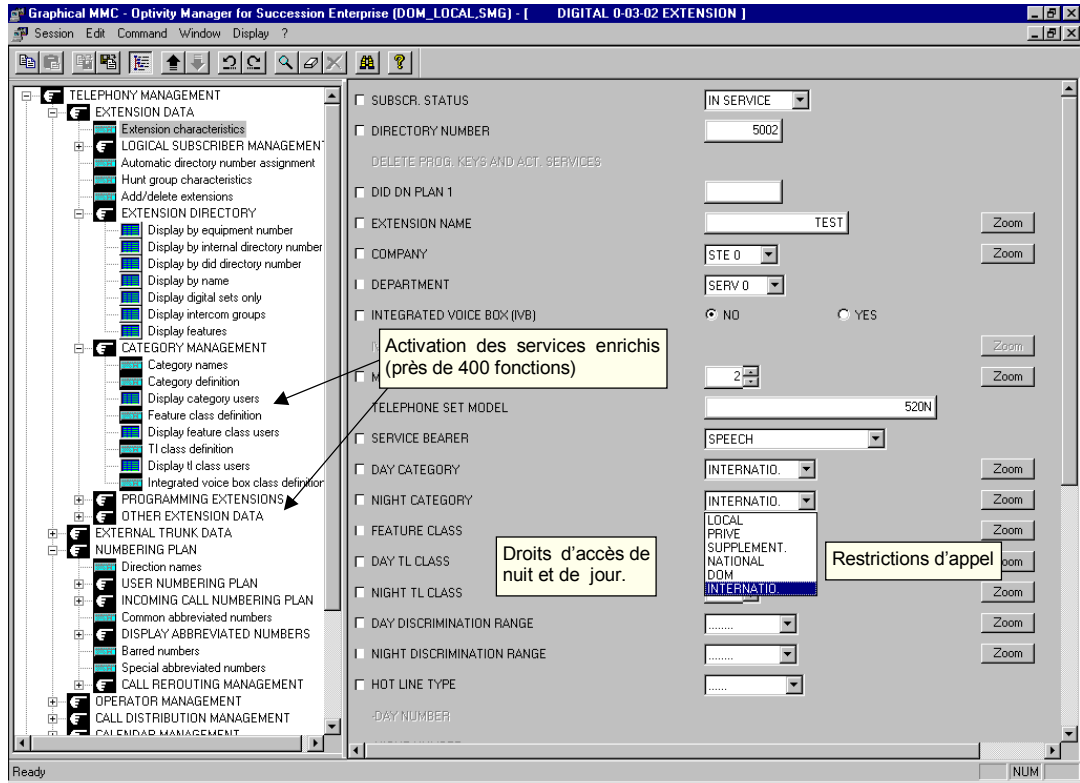
Subscriber number: 5090
 Subscriber type: IP subscriber - H323
 Subscriber number: DECT
 Subscriber number: DAS
 Subscriber number: CT2

Annotations:
 - ...ainsi que les services enrichis auxquels il a droit.
 - Fiche technique décrivant les caractéristiques du terminal... (points to the left sidebar)
 - Sur une passerelle, le terminal peut être IP ou non IP.
 - La messagerie vocale est gérée par la passerelle ou un serveur dédié.

Buttons: Confirm, Number range management

Sous l'arborescence *Subscription* (abonnement), nous avons créé l'abonné Montagnier, puis nous lui avons affecté un numéro de téléphone (le 5090) et spécifié le type de terminal qu'il utilise : dans son cas, il s'agit d'un téléphone IP (*IP subscriber – Etherset*). Nous lui avons

ensuite associé une fiche technique décrivant les caractéristiques du terminal (ici, *IP subscriber – H323*) et les services auquel il a droit, par exemple une messagerie vocale (*Voice mailbox*).



Parmi les services enrichis que le PABX peut gérer pour le terminal on trouve :

- l'identification de l'appelant ;
- la restriction d'appel (local, province, étranger, etc.), restrictions horaires ;
- le renvoi, transfert, filtrage ;
- le rappel sur occupation ;
- les numéros programmés (urgence, SVP) ;
- la gestion du second appel : parage, conférence, va-et-vient ;
- la messagerie vocale ;
- etc.

Assurer la qualité de service

Nous devons également penser à gérer la qualité de service sur l'ensemble de nos routeurs, en activant, par exemple, RSVP et les files d'attente appropriées.

La réservation de la bande passante doit être effectuée à la source des flux VoIP, c'est-à-dire au niveau des passerelles :

```
req-qos guaranteed-delay
```

← Demande une garantie sur la bande passante et sur le délai de transit

Ensuite, RSVP doit être activé sur toutes les interfaces de nos routeurs (comme indiqué au chapitre 14) :

```
ip rsvp bandwidth 45 15
fair-queue 64 256 3
```

45 kbit/s dédiés à RSVP, dont 15 kbit/s par flux (8 kbit/s pour le codec + overhead + signalisation)

3 files d'attentes pour RSVP pour 3 communications simultanées

Les débits à réserver dépendent des codecs choisis pour coder la voix (voir chapitre 12). Celui en question est le g729 qui nécessite 8 Kbit/s de bande passante.

Rappelons que les commutateurs sur nos LAN et que les files d'attente WFQ sur les routeurs prennent en compte la priorité des paquets IP. Il est donc intéressant d'affecter une valeur au champ "IP precedence", ce qui doit être fait à la source des flux VoIP, c'est-à-dire au niveau des passerelles :

```
#gwpar1
dial-peer voice 44 voip
ip precedence 5
```

← Priorité 5 = Critical (voir chapitre 13).

Par ailleurs, nos commutateurs doivent assurer la correspondance entre la priorité du paquet IP et celle affectée à la trame Ethernet, ce qui est réalisé comme suit :

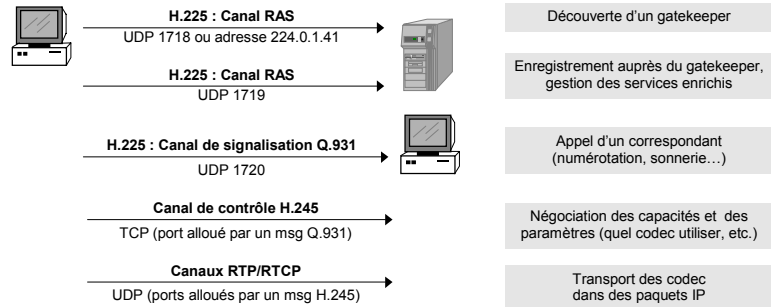
```
set qos acl ip politique_voip trust-ipprec
set qos ipprec-dscp-map 0 8 16 24 32 40 48 56
set qos dscp-cos-map 40:6
```

Transporter les flux multimédias

Un terminal H.323 utilise les services de la couche d'adaptation H.225 pour envoyer et recevoir des messages. En plus des signalisations Q.931 et RAS, cette couche réalise l'interface entre la gestion des canaux logiques par le terminal et la gestion des paquets sur un réseau IP. En fonction de leur nature, H.225 sélectionne ainsi le type de paquet dans lequel envoyer les données : TCP, UDP, RTP ou RTCP.

Canal logique	Protocole de transport / Numéro de port
Canal RAS : découverte d'un gatekeeper	UDP 1718 ou adresse multicast 224.0.1.41
Canal RAS : enregistrement auprès d'un gatekeeper	UDP 1719
Canal de signalisation Q.931	TCP 1720
Canal de contrôle H.245	Port TCP alloué dynamiquement par Q.931 (messages <i>Connect</i> , <i>Alerting</i> et <i>Call proceeding</i>)
Sessions RTP/RTCP	Port UDP alloué dynamiquement par H.245 (message <i>Open Logical Channel</i>)
T.120	TCP 1503

Figure 15-7.
*Interaction
des protocoles H.323.*



Le transport des flux multimédias est donc assuré par H.225 qui reprend les spécifications des RFC suivantes.

RFC	Sujet traité
1889	Spécifications de RTP et de RTCP
1890	Définition des profils pour les conférences audio et vidéo
2032	Transport des codecs vidéo H.261
2190	Transport des codecs vidéo H.263
2198	Transport des codecs audio

Le transport des flux audio et vidéo via RTP et RTCP

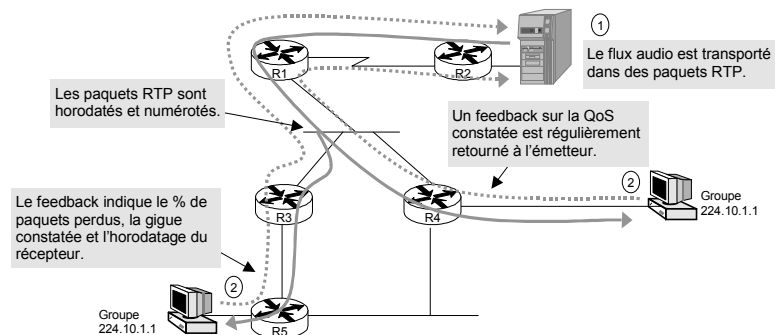
Une conversation téléphonique génère un flux audio qui est découpé en tranches représentant quelques millisecondes d'un son. De même, une image vidéo est découpée en tranches représentant un groupe de pixels. Une tranche correspond ainsi à un échantillonnage (*sampling*) d'un son ou d'une image vidéo numérisés et donc représentés par des octets (voir chapitre 12).

Ces octets sont placés dans des paquets **RTP** (*Real-time Transport Protocol*). Ce standard ne décrit aucun mécanisme de contrôle ni de récupération d'erreur : il se contente de définir le format des paquets et des données transportés. Ces paquets RTP sont ensuite encapsulés dans des paquets UDP, eux-mêmes transportés dans des paquets IP.

RTP est associé à **RTCP** (*RTP Control Protocol*), également transporté dans des paquets UDP, qui renvoie à l'émetteur un retour sur la qualité de service perçue par les récepteurs d'un flux.

Deux couples de ports UDP (source et destination) sont alloués pour un flux audio ou vidéo : l'un pour RTP, l'autre pour RTCP. Le protocole UDP étant orienté sans connexion, les sessions sont gérées par H.245 qui fait l'association avec ses numéros de canaux logiques.

Figure 15-8.
Rôles de RTP et de RTCP.



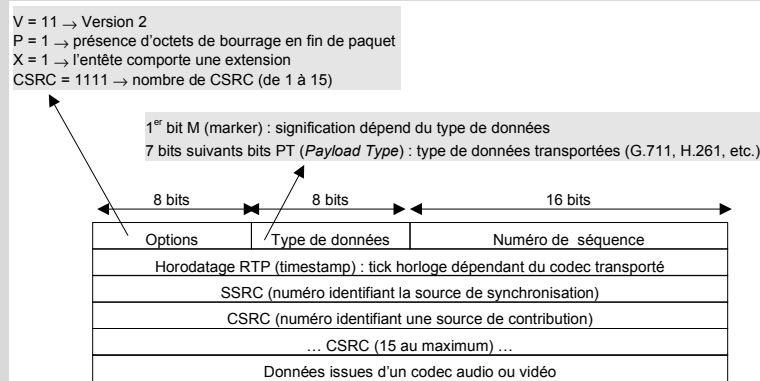
Un flux ne transporte jamais de données mixtes audio et vidéo. Par exemple, pour une visioconférence, les paquets audio et vidéo sont transportés dans deux sessions RTP différentes, chacune étant associée à une session RTCP.

Le codage peut être modifié au cours d'une session pour s'adapter à un changement de la qualité de transmission (débit, erreurs, etc.). L'émetteur peut également modifier la nature du flux au cours d'une session et basculer, par exemple, de la vidéo à l'audio.

LE POINT SUR RTP (RFC 1889 ET 1890)

Le protocole RTP (*Real-time Transport Protocol*) décrit le format des paquets transportant des flux audio ou vidéo.

Utilisant les services de la couche UDP, RTP n'assure pas pour autant un quelconque contrôle de flux, de reprise sur erreur ou même de contrôle d'intégrité du paquet. Son but est simplement de transporter quelques millisecondes de voix ou une portion d'image en y incluant des informations relatives au temps (synchronisation) et permettant d'identifier les émetteurs qui génèrent le signal audio ou vidéo.



Le numéro de port UDP affecté à RTP est alloué aléatoirement par l'application. Il doit simplement être pair. Le numéro suivant (forcément impair) est alors affecté à la session RTCP associée (voir l'encart suivant).

L'**horodatage**, qui dépend de l'horloge de l'émetteur, représente l'instant où le premier octet des données transportées a été échantillonné.

Le champ **SSRC** (*Synchronization Source*) est un numéro unique, généré de manière aléatoire, qui identifie l'émetteur du flux. Un mécanisme permet de gérer les problèmes de collision (deux SSRC identiques).

La RFC prévoit l'utilisation de **mixers** dont le rôle est de mélanger les flux pour n'en faire qu'un qui sera redistribué aux participants d'une conférence audio, par exemple. Il synchronise les flux combinés issus des différentes sources n'utilisant pas le même codage et ne disposant pas de la même source d'horloge ni de la même fréquence d'échantillonnage. Le paquet RTP résultant contient la liste des SSRC qui ont contribué à la formation du flux mélangé (champs **CSRC**, *Contributing Source*), par exemple, la liste de ceux qui ont parlé en même temps.

Le mixer s'apparente aux fonctions du MCU H.323, mais au niveau RTP, alors que ce dernier assure des fonctions plus évoluées liées à la signalisation. En pratique, un serveur de conférence implémente les fonctions de MCU et de mixer.

La RFC prévoit également l'utilisation de **translators** dont le rôle est de convertir des flux sans les mélanger.

Ce logiciel peut :

- convertir un codec en un autre ;
- convertir un flux multicast en un flux unicast ;
- convertir les ports UDP aléatoires en ports UDP fixes pouvant être filtrés par un firewall ;
- créer un tunnel sécurisé en chiffrant les données ;

Dans ces deux derniers cas, un translator s'utilise par paire.

LE POINT SUR RTCP (RFC 1889 ET 1890)

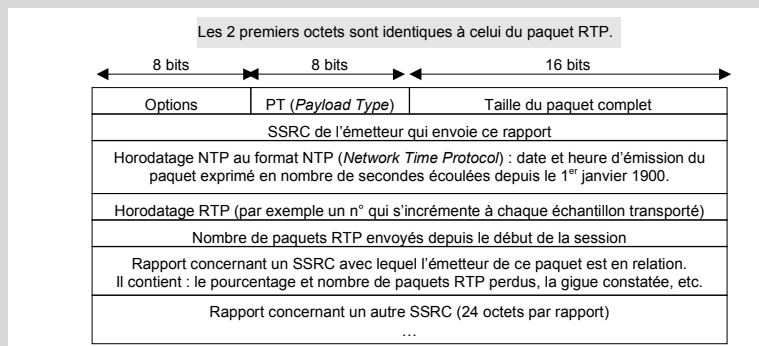
RTCP (*RTP Control Protocol*) décrit le format des informations remontées par le ou les récepteurs d'un flux. Il est utilisé concurremment avec RTP pour remplir trois fonctions :

- fournir un retour sur la qualité de service perçue par les récepteurs ;
- permettre d'identifier tous les participants à la conférence par un nom unique (un nom DNS, tel que *machine.laposte.fr*). Cela permet de retrouver le participant si le SSRC change en cours de session (suite à une réinitialisation, ou à la détection d'un conflit d'attribution de SSRC) ;
- permettre à chaque participant d'évaluer le nombre total de participants et d'adapter les flux en conséquence (fréquence d'émission des paquets).

Le paquet RTCP contient ainsi des informations relatives à l'**horodatage** des paquets reçus, la **gigue** (*jitter*) constatée par les récepteurs, le nombre de **paquets perdus** par période et en cumulé. Les informations recueillies par les émetteurs et les récepteurs permettent de vérifier que la qualité de service correspond à celle demandée via **RSVP**.

Il existe plusieurs types de paquets RTCP : rapport de récepteur (type RR – *Receiver Report*), rapport d'émetteur (type SR – *Sender Report*) si le récepteur est également un émetteur, description de la source (type SDES – *Source Description*), fin d'une participation (BYE) et fonctions spécifiques à une application (type APP), par exemple audio, vidéo, etc.

Ainsi, les rapports RR (PT = 201) et SR (PT = 200) prennent la forme suivante :



Le **rapport SEDS** contient, quant à lui, toutes les informations concernant l'émetteur : nom, adresse e-mail, numéro de téléphone, localisation géographique, nom du logiciel utilisé, etc.

Pour consommer le moins de bande passante possible :

- les différents types de paquets RTCP sont regroupés dans un seul paquet UDP ;
- la périodicité d'émission des paquets RTCP est aléatoire (entre 2 et 5 minutes) ;
- l'envoi des paquets RTCP par tous les participants doit être espacé de 2 à 7 secondes afin d'éviter les *bursts* ;
- le flux RTCP ne doit pas dépasser 5 % de la bande passante utilisée par RTP ;
- un participant ne doit pas générer plus de 20 % du total des flux RTCP.

Optimiser les flux multimédias

Un paquet RTP transporte 2 à 30 ms de voix, ce qui représente 20 à 150 octets selon le codec utilisé. La RFC 2298 montre l'exemple d'un paquet RTP transportant 20 ms de voix échantillonnée à 8 KHz et dont les données utiles ne représentent que 98 octets.

Codec	Données numériques
G.722	8 bits par échantillonnage
G.723	30 ms par trame
G.728	3,1 ms par trame
G.729	10 ms par trame

Afin de réduire l'impact de l'overhead, on pourrait augmenter la taille des données transportées en envoyant plusieurs trames dans un seul paquet. L'inconvénient de cette approche est qu'elle engendre un délai de transit supplémentaire dans une passerelle ou au départ du PC émetteur. Par exemple, au lieu d'attendre 30 ms pour envoyer une trame G.723, on attendra 60 ms pour envoyer deux trames. Résultat : le délai d'attente de la première trame du paquet est allongé de 30 ms, et ce pour une trame sur deux dans le flux.

Les techniques suivantes sont pour cela privilégiées.

Compression des en-têtes

La taille minimale d'un paquet RTP (sans la liste des contributeurs) est de 12 octets, auxquels il faut ajouter 8 octets pour UDP et 20 pour IP, soit 40 octets en tout.

Afin de diminuer cet overhead, les routeurs permettent de compresser ces en-têtes en réduisant leur taille cumulée à 2 ou 5 octets sur les interfaces WAN. Par exemple, les commandes suivantes permettent de compresser les en-têtes de 16 sessions RTP simultanées :

```
int s 0
encapsulation ppp
ip rtp header-compression
ip rtp compression connections 16
```

Le principe repose sur le fait que l'en-tête varie très peu d'un paquet à l'autre : seuls les numéros de séquence et de contrôle d'erreur changent. L'activation de la compression consiste alors à n'envoyer que les données qui ont changé d'un en-tête à l'autre (RTP/UDP/IP), ce qui représente entre 2 et 5 octets la plupart du temps. Cette fonction est surtout efficace pour les liaisons inférieures à 2 Mbit/s.

Utilisation des mixers

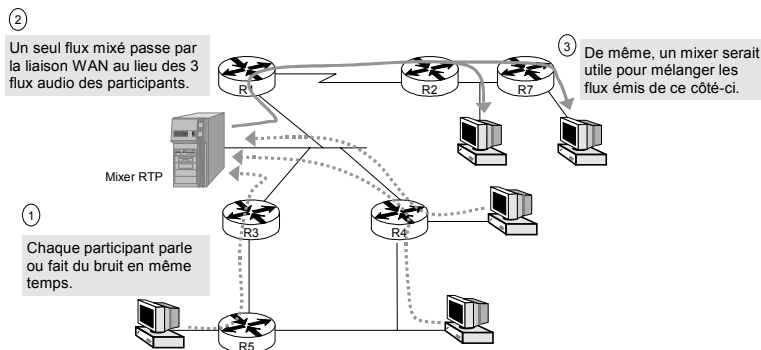
Le son perçu par un participant est la résultante des voix et des bruits émis par les autres. Pourtant, celui-ci reçoit N flux audio, provenant des N autres participants. Son PC se contente de restituer un son composite.

Afin de diminuer le débit généré par ces flux sur notre réseau, et surtout sur les liaisons WAN limitées en bande passante, il serait intéressant de créer un son composite en amont.

Le **mixer RTP** répond à ce besoin puisqu'il permet de fusionner les différents flux audio en un seul. Rappelons, par ailleurs, qu'un mixer permet également de convertir les codecs audio entre participants ne disposant pas des mêmes équipements.

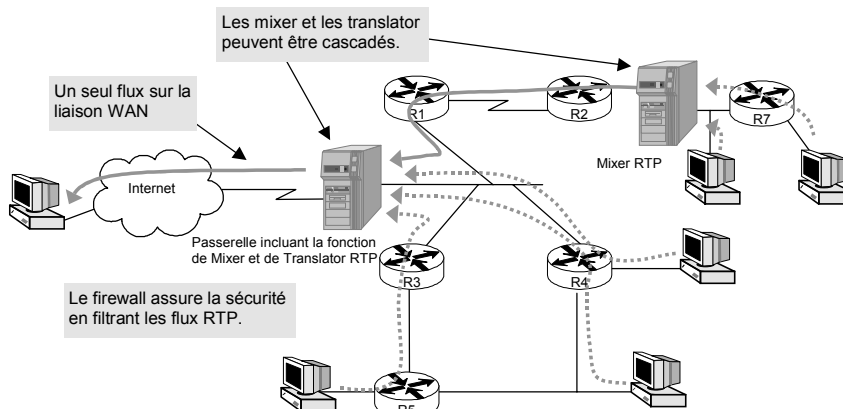
Dans notre réseau, l'installation d'un mixer de chaque côté de la liaison WAN permet de diminuer considérablement les flux générés par une conférence audio. De plus, alors que les participants peuvent envoyer leur flux dans des paquets unicast, le flux envoyé par le mixer pourra utiliser des paquets multicast afin de réduire encore davantage la charge globale du réseau.

Figure 15-9.
Utilisation d'un mixer pour optimiser les flux.



L'intranet est connecté à l'Internet par une liaison (spécialisée, Frame Relay ou autre) dont le débit est souvent limité. Si plusieurs participants à une même conférence sont répartis entre notre intranet et l'Internet, là encore, l'utilisation d'un mixer permet de réduire la charge de la liaison WAN.

Figure 15-10.
Intégration d'un mixer et d'un translator dans un firewall.

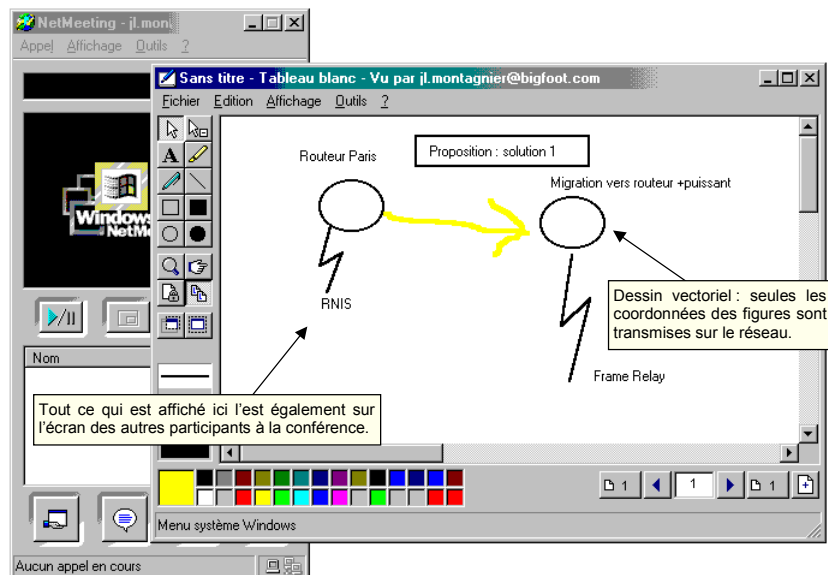


Échanger des données multimédias

Les participants à une visioconférence peuvent vouloir recevoir un document que leur remet l'animateur, par exemple. Le standard T.120 de l'ITU-T offre à ces participants la possibilité d'utiliser quatre applications :

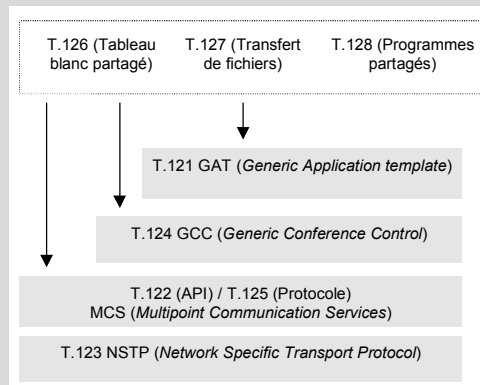
- transfert de fichiers multipoint ;
- partage d'un presse-papiers commun ;
- partage d'un tableau blanc virtuel (*whiteboard sharing*) qui permet à chaque participant de dessiner sur le même schéma ;
- déport écran/clavier pour une prise de contrôle à distance et une démonstration à tous les participants.

On retrouve ici des applications classiques, mais sous la forme du partage par plusieurs utilisateurs.



LES DONNÉES AU SEIN D'UNE CONFÉRENCE MULTIMÉDIA (ITU-T T.120)

La norme T.120 décrit un cadre fonctionnel général permettant à des utilisateurs de **partager des données au sein d'une conférence** *via* des sessions TCP gérées par H.245.



Les protocoles utilisés sont les suivants :

- **GAT** (*Generic Application Template*), norme **T.121**. Définit les interfaces d'accès (API) pour des applications T.120 : entrées et sorties des conférences, négociations des fonctionnalités, etc.
- **MCS** (*Multipoint Communication Service*), normes **T.122** pour l'interface d'accès et **T.125** pour le protocole. Assure la diffusion des données au sein d'une conférence selon trois topologies : en étoile autour d'un gestionnaire central (*top provider*), en cascade autour de plusieurs gestionnaires et d'un *top provider*, ou chaînage (*daisy chain*).
- **NSTP** (*Network Specific Transport Protocol*), norme **T.123**. Assure l'interface entre les applications T.120 et le réseau. Il réalise l'adaptation à différents supports, tels que RNIS et TCP/IP, tout en gérant les erreurs de transmission.
- **GCC** (*Generic Conference Control*), norme **T.124**. Il s'agit du logiciel qui organise la conférence. Il gère la liste des participants, les accepte ou les refuse selon les mots de passe et les droits d'accès, décide à qui passer la main, assure la cohérence des informations échangées (mise à jour en temps réel et simultanément pour tous les participants) et surveille les ressources utilisées par le MCS.

QUATRIÈME PARTIE

Gérer son réseau

16

Administrer son réseau IP

Votre réseau d'entreprise est désormais opérationnel. Il s'étend sur plusieurs sites. Cependant, les problèmes vous guettent. En effet, plus le réseau est important, plus la probabilité qu'une panne survienne à un endroit ou à un autre est importante. C'est statistique.

En outre, plus un réseau est important, plus il est difficile à gérer. Il convient donc d'utiliser des outils qui simplifient sa gestion et diminuent donc le nombre potentiel de pannes.

Dans ce chapitre vous apprendrez :

- à utiliser les outils de base pour le diagnostic réseau ;
- à installer un serveur DHCP qui vous facilitera la vie.

Les utilitaires de base

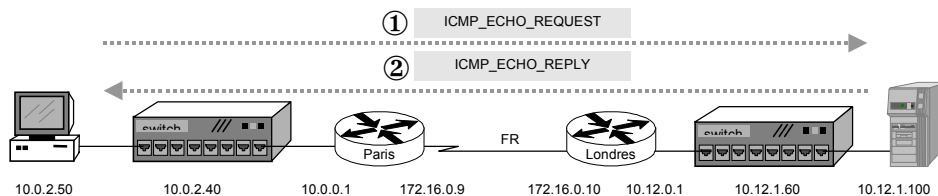
Le ping

Deux programmes doivent faire partie du « kit de survie » de tout administrateur réseau : ping et Traceroute.

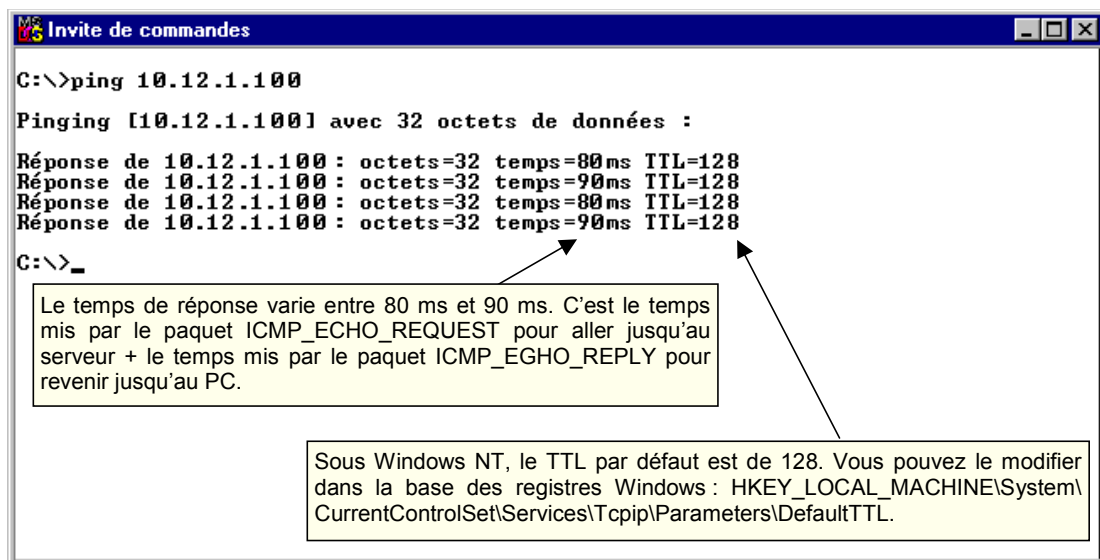
La commande de base est le **ping** (*Packet Internet Groper*, ou *ping-pong*). Ce programme permet de savoir si une station IP est active et, plus généralement, si le réseau fonctionne correctement entre deux points, par exemple, entre votre PC à Paris et le serveur de Londres. Autre fonction intéressante, le programme donne également le temps de réponse mesuré.

Figure 16-1.

Le ping.



La commande ping s'appuie sur le protocole ICMP (*Internet Control Message Protocol*) qui fait partie de la couche IP. Ce protocole est utilisé pour toutes les opérations qui ont trait à la gestion du réseau IP, et ce, de façon transparente pour les utilisateurs.



Le traceroute

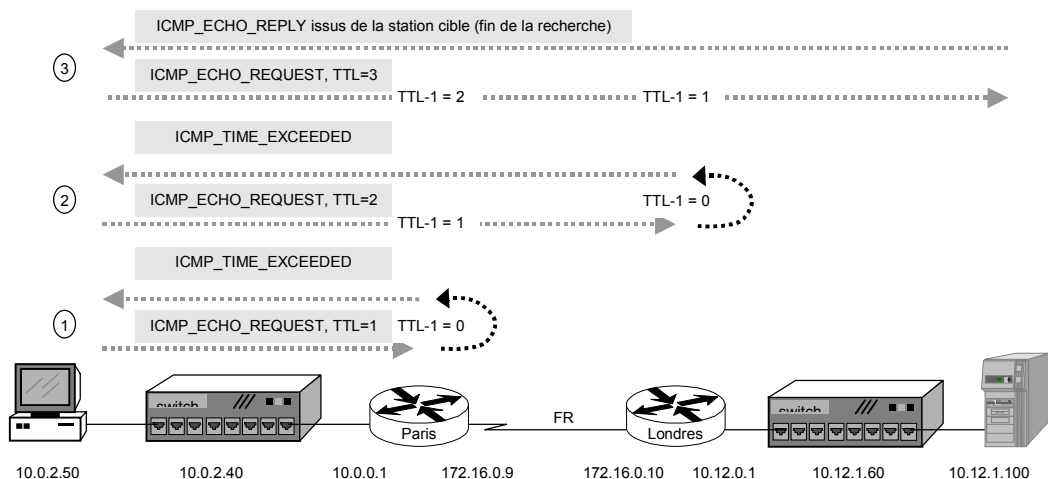
L'autre commande de base est le **traceroute**. Ce programme utilise des mécanismes propres à IP et ICMP pour afficher à l'écran la route empruntée par un paquet IP en plus du temps de réponse. Son principe de fonctionnement est le suivant :

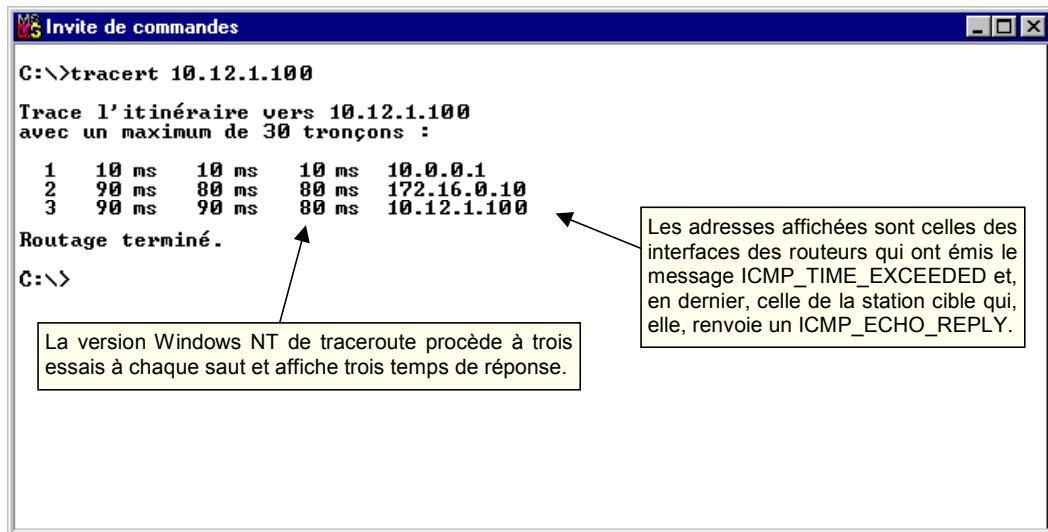
- Le programme envoie un paquet ICMP_echo_request à destination de la machine cible avec un TTL (*Time To Live*) égal à 1.
- Le premier routeur reçoit ce paquet, décrémente le TTL de 1, et constate qu'il est égal à 0. Il détruit le paquet, et renvoie à l'émetteur un message ICMP_Time_exceeded.
- Le programme enregistre l'adresse IP du routeur qui a envoyé de ce message ainsi que le temps écoulé depuis l'émission du paquet ICMP_Echo_request.
- Le programme continue de même en incrémentant le TTL de 1 à chaque paquet ICMP_Echo_request émis. Le paquet ira donc un saut plus loin que le précédent, et le routeur suivant répondra.

Le mécanisme du TTL (*Time To Live*) est expliqué dans l'encart « Le point sur IP v4 », au chapitre 7.

Certaines implémentations de traceroute utilisent un paquet UDP sur un port quelconque à la place d'un paquet ICMP_Echo_request. La RFC 1393 (statut expérimental) propose, quant à elle, un autre algorithme qui repose sur un message ICMP_Traceroute (type 30). Il est cependant rarement implémenté, aussi bien sur les stations que sur les routeurs.

Figure 16-2.
Le traceroute.



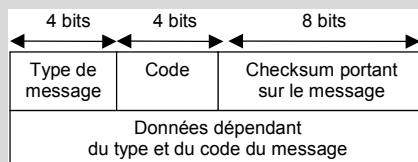


La commande ping sera plutôt utilisée pour savoir si un nœud IP est actif et joignable ainsi que le temps de réponse de bout en bout. Elle offre en outre davantage de possibilités de paramétrage (taille et nombre des paquets, enregistrement de la route, etc.). La commande traceroute permet, quant à elle, de savoir quelle route est empruntée (par exemple, le chemin principal ou celui de backup) et quelles parties du réseau engendrent les temps de réponse les plus longs.

De nombreux utilitaires de ce type sont disponibles sur les sites ftp.lip6.fr, tucows.club-internet.fr et www.winsite.com.

LE POINT SUR ICMP (RFC 792, 950 ET 1256)

ICMP (*Internet Control Message Protocol*) regroupe près de trente types de messages permettant aux nœuds d'un réseau d'échanger des informations de gestion relatives à IP.



Excepté pour les messages Echo des Timestamp, le paquet ICMP contient une copie partielle du paquet original ayant causé l'erreur (en-tête IP + 8 premiers octets des données).

Le tableau de la page suivante recense les types et codes existants à ce jour. Lorsque cela n'est pas précisé, les messages sont générés par les routeurs et les stations IP. Sauf indication contraire, les messages ICMP sont définis dans la RFC 792.

...

LE POINT SUR ICMP (RFC 792, 950 ET 1256 – SUITE)

Type	Message	Code et description
3	Destination Unreachable	0 <i>network unreachable</i> : le routeur ne connaît pas la route 1 <i>host unreachable</i> : le routeur ne peut pas trouver la station 2 <i>protocol unreachable</i> : le protocole demandé n'est pas actif. 3 <i>port unreachable</i> : aucun programme ne répond sur ce port TCP ou UDP. 4 <i>fragmentation needed and DF set</i> : le routeur a reçu un fragment alors que la fragmentation est interdite (bit DF du paquet IP positionné à 1). 5 <i>source route failed</i>
11	Time exceeded	0 Si un routeur reçoit un paquet avec un TTL à 0, il envoie ce message à l'émetteur. 1 Si une station n'obtient pas tous les fragments d'un message dans le temps imparti, elle envoie ce message à l'émetteur.
12	Parameter Problem	0 Des paramètres incorrects ou inconsistants dans l'en-tête du paquet IP ont été détectés (un pointeur indique la position de l'erreur dans l'entête)
4	Source Quench	0 Le routeur, ou la station, est congestionné (ou régule le trafic selon un algorithme propre) et demande à l'émetteur de réduire son flux. Les paquets en excès peuvent être détruits.
5	Redirect	0 <i>for the Network</i> : le routeur a détecté une meilleure route et indique à la station quel routeur solliciter 1 <i>for the Host</i> : idem pour une station 2 <i>for the TOS and Network</i> : idem avec le champ TOS correspondant. 3 <i>for the TOS and Host</i> : idem pour une station.
8	Echo request	0 Demande au récepteur de renvoyer un Echo reply (un identifiant et un numéro de séquence identifient le message).
0	Echo reply	0 Réponse à un Echo request. Les données contenues dans le message Echo request doivent être reportées dans ce message.
13	Timestamp	0 Indique le nombre de millisecondes écoulé depuis 00h00 GMT lorsque le message Timestamp a été envoyé. Utilisé pour évaluer le temps de transit.
14	Timestamp reply	0 Indique le nombre de millisecondes écoulé depuis 00h00 GMT lorsque le message Timestamp a été reçu, ainsi que la valeur de ce nombre lorsque la réponse a été envoyée.
15	Information request	0 (obsolète) Permettait aux stations d'obtenir leur adresse IP. Mécanisme remplacé par les protocoles RARP, puis BOOTP et DHCP.
16	Information reply	0 Idem (obsolète).
17	Mask request	0 (RFC 950) Permet à une station d'obtenir le masque IP de son sous-réseau.
18	Mask reply	0 (RFC 950) Réponse du routeur à une demande de masque.

...

LE POINT SUR ICMP (RFC 792, 950 ET 1256 – FIN)

Type	Message	Code et description
9	Router Advertisement	0 (RFC 1256) Émis périodiquement par un routeur pour indiquer l'adresse IP de son interface. Permet aux routeurs de découvrir leurs voisins, et aux stations de découvrir leur passerelle par défaut.
10	Router Sollicitation	0 (RFC 1256) Lors de son initialisation, une station peut demander à un routeur de s'annoncer immédiatement. Les routeurs n'envoient, en principe, aucune sollicitation, mais attendent les annonces.

L'IANA recense d'autres messages ICMP, soit expérimentaux, soit pour IPv6.

Observer ce qu'il se passe sur son réseau

Si votre réseau présente un dysfonctionnement et que, malgré toutes vos investigations, vous n'avez pas trouvé d'où provient le problème, il ne vous reste plus qu'à l'ausculter, c'est-à-dire observer les données qui y circulent.

Même lorsque le réseau semble bien fonctionner, il n'est pas inutile d'y jeter un coup d'œil, car bien souvent des erreurs (collision, paquets corrompus, flux non identifié, trafic censé ne pas être présent sur ce segment, etc.) se produisent. Ces erreurs ne sont alors pas perceptibles, mais peuvent le devenir sous certaines conditions, par exemple lorsque la charge réseau augmente. Une **maintenance préventive** permet donc d'éviter le pire.

L'**analyseur réseau** est l'outil tout indiqué pour ce type de situation. Il permet :

- de capturer toutes les trames qui circulent sur un segment Ethernet ;
- d'analyser le contenu de toutes les couches réseau, de la trame aux données applicatives en passant par le paquet IP ;
- de déterminer si des erreurs se produisent (collision, erreur de transmission, etc.) et en quelle proportion ;
- de connaître les temps de réponse précis (au millième de seconde près).

En positionnant des filtres, il est possible de suivre précisément les échanges entre deux stations, soit à partir de leurs adresses MAC, soit à partir de leurs adresses IP.

COMMENT UN ANALYSEUR RÉSEAU FONCTIONNE-T-IL ?

Un analyseur réseau est un logiciel capable de décoder idéalement tous les protocoles existants, du niveau 2 au niveau session. Il fonctionne de concert avec une carte réseau, de préférence haut de gamme, capable de capturer toutes les trames, même à pleine charge.

Le coût d'un tel produit dépend donc du nombre de protocoles reconnus et de la carte d'acquisition : Ethernet, Token-Ring ou ATM pour les LAN, et série synchrone pour les liaisons WAN en Frame Relay, ATM, etc.

La carte réseau doit fonctionner en mode **promiscus**. Dans un mode de fonctionnement normal, une carte ne prend en compte que les trames multicast et de broadcast, ainsi que celles dont l'adresse de destination MAC correspond à celle qui est programmée dans sa mémoire (PROM, Flash, etc.). En mode promiscus, la carte prend en compte toutes les trames. Toutes les cartes réseau ne supportent pas le mode promiscus.

Sniffer Pro - Local/3Com Etherlink III (3c509) ISA Adapter_4 (Line speed at 10 Mbps) - [Snif1.cap : 41/54 Ethernet frames]

La couche liaison (DCL = Data Link Control) correspond ici à une trame Ethernet.

No.	Status	Source Address	Dest Address	Summary	Len	Rel.
37		NT001	[10.255.255.255]	BROWSER: Local Master NT001 Announce	243	0:(
38		JLM	[10.255.255.255]	BROWSER: Announce Host JLM	256	0:(
39		Fujitsu696C42	Broadcast	ARP: C PA=[10.0.0.100]NT001 PRO=IP	60	0:(
40		10040B5646E9	Fujitsu696C42	ARP: R PA=[10.0.0.100]NT001 HA=10040B5646E9	60	0:(
41		JLM	NT001	ICMP: Echo	74	0:(
42		NT001	JLM	ICMP: Echo reply	74	0:(
43		JLM	NT001	ICMP: Echo	74	0:(

Détail du paquet

Requête ARP suivie de sa réponse, qui permet d'envoyer la trame Ethernet à l'adresse MAC correspondant à l'adresse IP à qui est envoyé le paquet ICMP_echo_request (ping).

Adresses MAC source et destination. L'analyseur a repéré qu'il s'agissait d'une carte Fujitsu et a remplacé la partie constructeur de l'adresse MAC par le nom.

La trame Ethernet transporte un paquet IP. Il s'agit d'une trame Ethernet v2, car la valeur du champ est supérieure à 1 500.

Données brutes telles qu'elles circulent sur le segment Ethernet.

Expert / Decode
For Help, press F1
Démarrer Sniffer Pro - Local/3...

L'analyseur réseau peut également être utilisé comme outil d'analyse pour :

- mesurer la volumétrie générée par une application entre un client et un serveur ;
- évaluer la part de responsabilité du réseau, du serveur et du client dans le temps de réponse global entre un client et un serveur ;
- surveiller la charge du réseau pendant 24 heures ou sur une semaine ;
- déterminer quelles stations génèrent le plus de trafic ;
- déterminer la répartition du trafic par protocole, par adresse IP, etc.

Enfin, l'analyseur réseau offre souvent des fonctions évoluées, telles que :

- une minuterie (*trigger*) qui permet de déclencher et d'arrêter la capture sur réception d'une trame particulière (adresse, données, etc.) ;
- un générateur de trafic pour vérifier le comportement du réseau et des applications à pleine charge (par exemple si trop d'erreurs surviennent à partir d'une certaine charge, le câblage est sans doute en cause) ;
- la possibilité de rejouer un échange de trames préalablement capturées.

Piloter son réseau

Si votre réseau prend de l'ampleur — le nombre des équipements (routeurs, concentrateurs, commutateurs, etc.) augmente, et ces derniers sont répartis sur différents sites —, il devient de plus en plus nécessaire de centraliser la gestion des équipements.

Pour ce faire, la famille des protocoles TCP/IP propose le protocole SNMP (*Simple Network Management Protocol*) qui permet l'échange d'information entre **une station d'administration** (le client) et des **agents** (les serveurs) implantés dans chaque équipement. On parle alors d'**agent SNMP** ; celui-ci se présente sous forme d'un petit programme qui répond aux requêtes SNMP émises par la station d'administration.

Quelle station d'administration ?

Une station d'administration, ou **plate-forme d'administration**, est constituée d'un ordinateur sous Windows NT/2000 ou sous Unix, ainsi que d'un logiciel, tel qu'OpenView de Hewlett-Packard, Tivoli d'IBM, Unicenter de Computer Associates, Spectrum de Cabletron, etc.

Ces logiciels haut de gamme (environ 100 000 francs) sont en fait des boîtes à outils sur lesquelles s'installent des **modules dédiés** à chaque constructeur (CiscoView pour les équipements Cisco, Optivity pour ceux de Bay Networks, etc.). Ces modules peuvent, par ailleurs, fonctionner de manière autonome.

L'intérêt d'une telle plate-forme est de fédérer la gestion d'un parc d'équipements hétérogène autour d'une gestion centralisée des alarmes et d'une carte réseau sur laquelle s'affichent les équipements découverts dynamiquement.

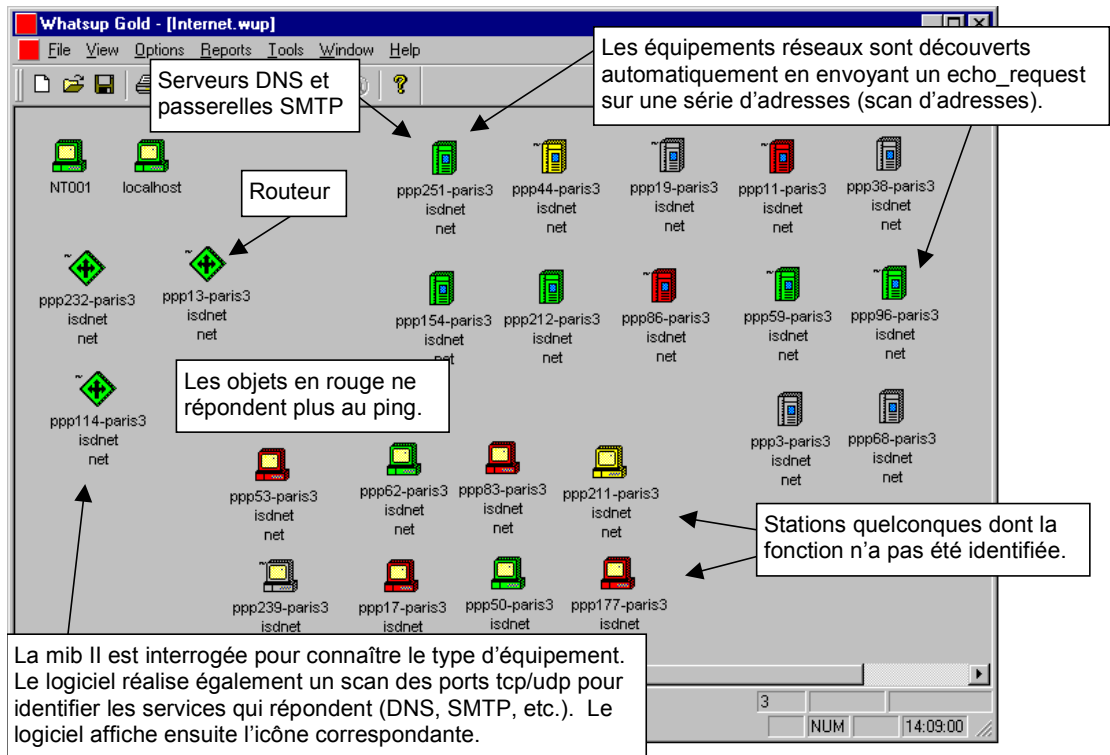
En fait, les plates-formes d'administration nécessitent un paramétrage très important dont le coût peut être plus élevé que celui du matériel et du logiciel réunis. Elles sont pour cela réservées aux grands réseaux, notamment chez les opérateurs.

Pour un réseau de plus petite taille, il est préférable d'utiliser les modules des constructeurs en mode autonome : si vous disposez d'un parc d'équipements homogène, nul besoin d'investir dans une « usine à gaz ». L'intérêt est de pouvoir visualiser graphiquement les équipements et de cliquer sur les cartes et ports que vous voulez configurer.

Certains administrateurs de grands réseaux se dispensent même de ce type de logiciel, préférant utiliser Telnet, TFTP et les fichiers de configuration en mode texte, le ping et le trace-route étant utilisés pour les dépannages quotidiens.

Pour quelle utilisation ?

Si vous désirez néanmoins visualiser et surveiller votre réseau de manière graphique, vous pouvez toujours utiliser des petits logiciels, tels que *Whatsup*. Le principe est identique à celui des plates-formes, mais avec un peu moins de fonctionnalités.



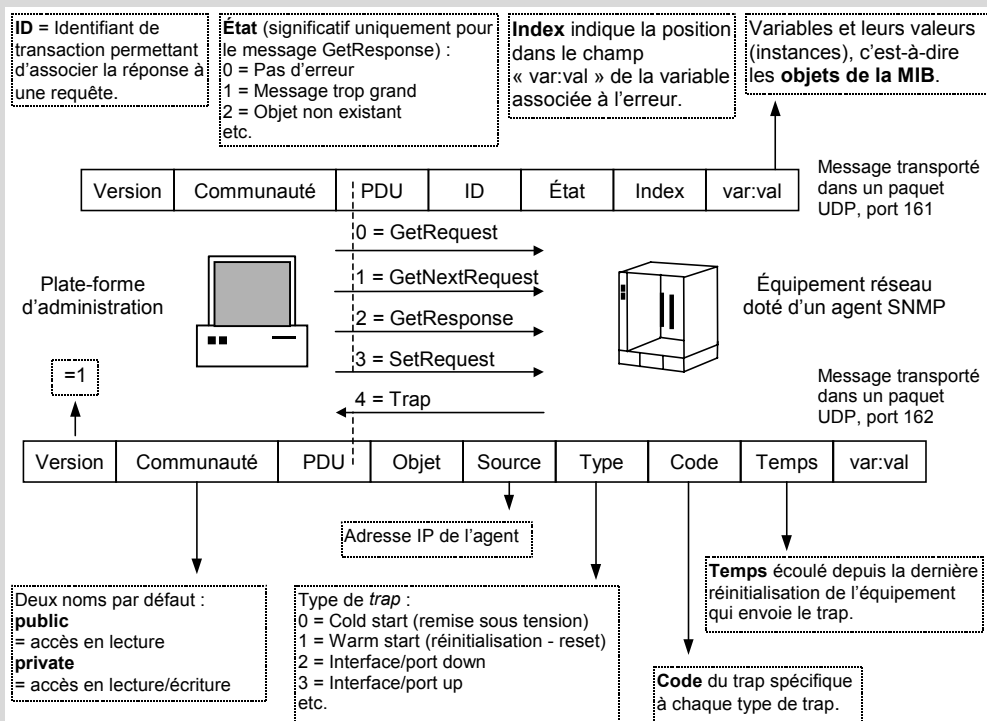
Cette carte a été obtenue en scannant les adresses IP d'un fournisseur d'accès à l'Internet (ISP). Le logiciel a ainsi trouvé des routeurs, des serveurs DNS, des passerelles de messagerie SMTP, ainsi qu'un certain nombre de stations non identifiées.

La première tâche est d'agencer les icônes qui apparaissent dans le désordre. L'administrateur peut ensuite dessiner un réseau et positionner les objets dessus, de manière à faire correspondre la carte à la réalité.

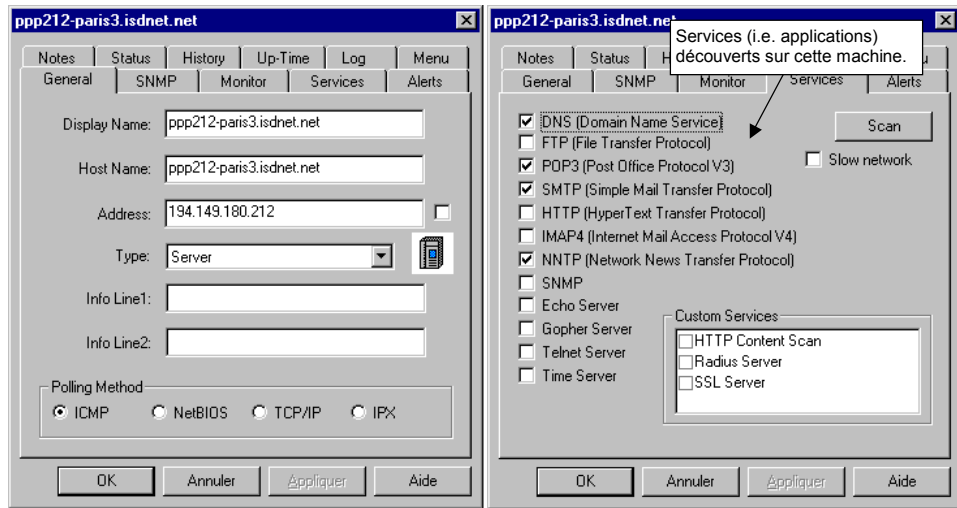
En sélectionnant une icône, il est alors possible d'opérer plusieurs actions sur l'équipement : ajouter des informations complémentaires, interroger son agent SNMP, surveiller des paramètres de cet agent, positionner des seuils d'alerte, etc.

LE POINT SUR SNMP v1 (RFC 1157, 2571, 2572)

Le protocole SNMP (*Simple Network Management Protocol*) est utilisé pour piloter tous les équipements du réseau (routeurs, commutateurs, concentrateurs, serveurs, etc.) à partir d'une **station d'administration**. Il est ainsi possible de **configurer** à distance les équipements (activation d'une interface, ajout d'une adresse IP, etc.) et de récupérer les paramètres actifs. Inversement, un équipement peut envoyer une alarme à la station d'administration via un **trap SNMP**.



Presque tous les équipements réseau intègrent un **agent SNMP**. Ce logiciel réalise l'interface entre les **requêtes SNMP** et la base de donnée **MIB** (*Management Information Base*) qui regroupe tous les paramètres de l'équipement.



Le problème de ce type de logiciel est qu'il faut sans cesse le mettre à jour, car le réseau ne cesse d'évoluer. L'autre problème tient à la gestion des alertes : le logiciel doit être très précisément paramétré pour ne générer que des alarmes réelles. Ces deux activités peuvent prendre beaucoup de temps à l'administrateur.

```
RFC1213-MIB DEFINITIONS ::= BEGIN
```

```
IMPORTS
```

```
    mgmt, NetworkAddress, IpAddress, Counter, Gauge, TimeTicks
```

```
FROM RFC1155-SMI
```

```
OBJECT-TYPE
```

```
FROM RFC-1212;
```

Importe ces groupes de la MIB SMI et tous les objets de la MIB-I.

Cette MIB définit le groupe MIB-II situé sous le groupe Management.

```
mib-2 OBJECT IDENTIFIER ::= { mgmt 1 }
```

```
DisplayString ::= OCTET STRING
```

```
PhysAddress ::= OCTET STRING
```

```
-- groups in MIB-II
```

```
    system OBJECT IDENTIFIER ::= { mib-2 1 }
```

```
    interfaces OBJECT IDENTIFIER ::= { mib-2 2 }
```

```
    at OBJECT IDENTIFIER ::= { mib-2 3 }
```

```
    ip OBJECT IDENTIFIER ::= { mib-2 4 }
```

```
    icmp OBJECT IDENTIFIER ::= { mib-2 5 }
```

```
    tcp OBJECT IDENTIFIER ::= { mib-2 6 }
```

```
    udp OBJECT IDENTIFIER ::= { mib-2 7 }
```

```
    egp OBJECT IDENTIFIER ::= { mib-2 8 }
```

```
-- cmot OBJECT IDENTIFIER ::= { mib-2 9 }
```

Définition de l'objet **system**, identifiant n° 1 dans le groupe **mib-2**

Noms et identifiants dans la MIB 2.


```

transmission OBJECT IDENTIFIER ::= { mib-2 10}
snmp OBJECT IDENTIFIER ::= { mib-2 11 }

-- the System group
sysDescr OBJECT-TYPE
    SYNTAX DisplayString (SIZE (0..255))
    ACCESS read-only
    STATUS mandatory
    DESCRIPTION
        "Par ex. : Cisco 761 Software Version c760-i..."
    ::= { system 1 }

sysObjectID OBJECT-TYPE
    SYNTAX OBJECT IDENTIFIER
    ACCESS read-only
    STATUS mandatory
    DESCRIPTION
        "Par Ex. : Cisco2503"
    ::= { system 2 }
    .....

ifNumber OBJECT-TYPE
    SYNTAX INTEGER
    ACCESS read-only
    STATUS mandatory
    DESCRIPTION "Nombre d'interfaces, par exemple, 3"
    ::= { interfaces 1 }
    .....
END

```

Description de l'objet **sysDescr**,
identifiant n° 1 dans le groupe
system.

Description de l'objet **sysObjectID**,
identifiant n° 2 dans le groupe
system.

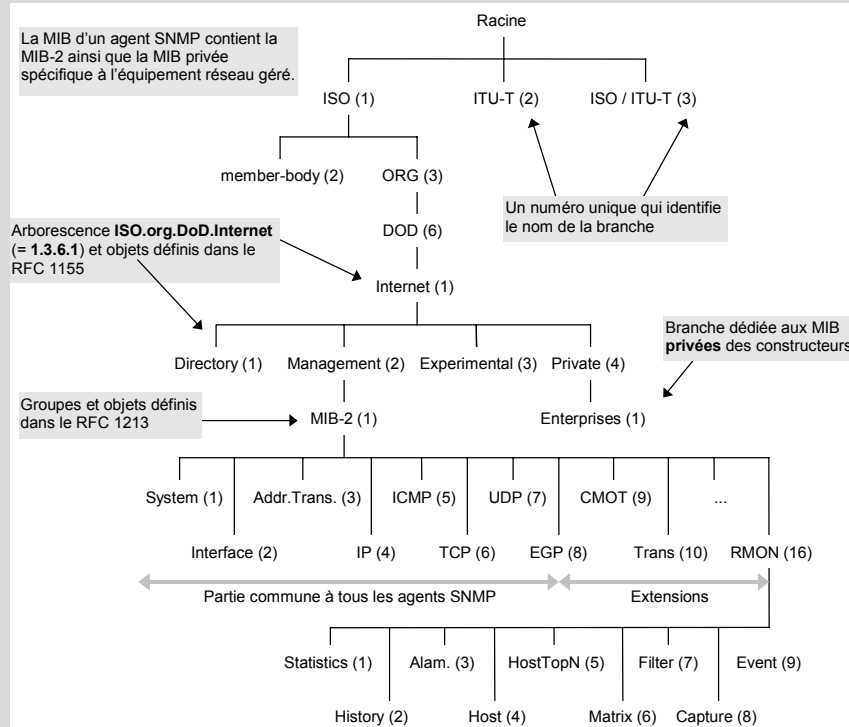
Description de l'objet **ifNumber**,
identifiant n° 1 dans le groupe
interfaces.

Pour que la station d'administration puisse interroger la MIB, il faut en premier lieu la **compiler** à partir du fichier en syntaxe ASN.1. Le fichier texte est alors intégré sous une autre forme (binaire généralement) dans le gestionnaire de MIB du logiciel d'administration.

Le moyen le plus simple de visualiser la MIB d'un équipement réseau est d'utiliser le module dédié, propre à chaque constructeur. La manipulation des variables est alors transparente, puisque le module affiche graphiquement l'équipement, par exemple, un commutateur. Il suffit alors de cliquer sur un port et de sélectionner les options qui vous sont proposées: activation/désactivation, vitesse (10, 100 ou *autosense*), nombre d'octets émis et reçus, etc.). Par ailleurs, l'interface graphique affiche les éléments dans différentes couleurs en fonction des alarmes (*trap*) qui lui sont remontées.

LE POINT SUR LA MIB (RFC 1212, 1213, 1155 ET 2863)

Les agents SNMP interagissent avec la **MIB** (*Management Information Base*) qui contient tous les paramètres de l'équipement réseau. Cette base de données se présente sous forme d'arborescence, normalisée ISO, dans laquelle une branche est réservée à l'Internet. Chaque objet de l'arborescence est identifié par un numéro.



La structure de cette base de données est décrite dans la syntaxe **ASN.1** (*Abstract Syntax Notation 1*) normalisée ISO 8824. Un fichier MIB comporte deux parties (RFC 1212), la première décrivant les types et groupes d'objets (macro **Definitions**), la seconde décrivant les objets (série de macro **Object-Type**).

La macro Definitions contient deux **clauses** :

Import : importe des définitions d'autres fichiers MIB.

Object Identifier : définit un nouveau groupe (nom + identifiant).

La macro Object-Type contient trois clauses obligatoires, qui prennent les valeurs suivantes :

Syntax = integer | object identifier | octet string | networkaddress | ipaddress.

Access = read-only | read-write | write-only | not-accessible.

Status = mandatory | optional | obsolete | deprecated.

Et quatre autres facultatives, qui prennent les valeurs suivantes :

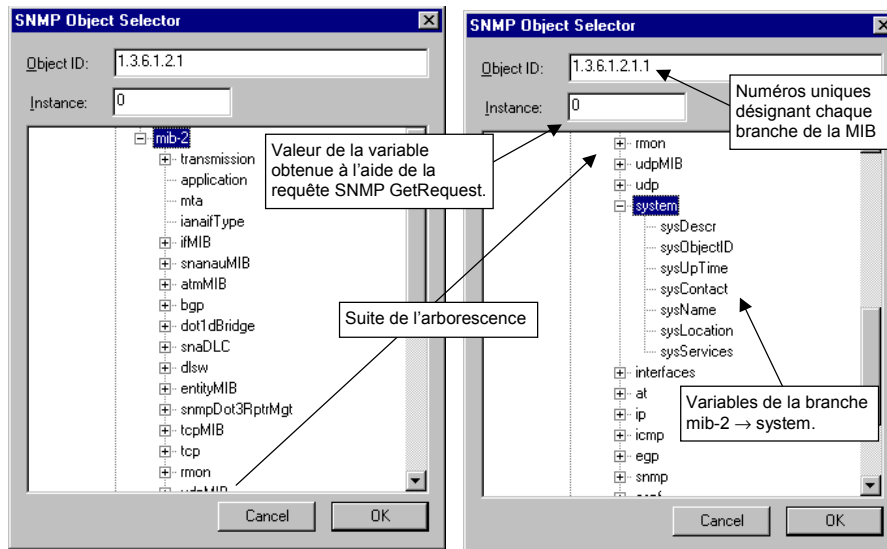
Description = "texte décrivant l'objet".

Reference = référence à un autre objet.

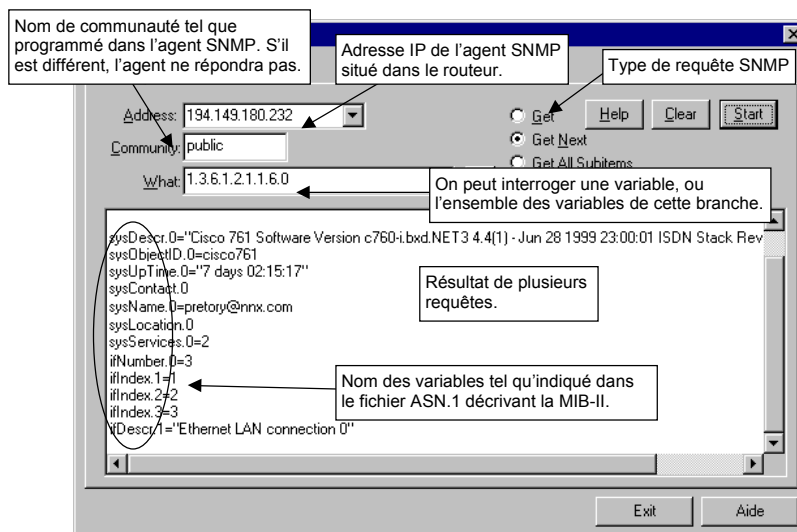
Index = noms d'objets dans un objet structuré, par exemple, "ifIndex" pour l'objet *ifEntry* qui contient une liste d'interfaces.

Defval = valeur par défaut de l'objet, par exemple, "sysDescr" pour une syntaxe *Object Identifier*, ou "1" pour une syntaxe *Integer*.

La seconde solution est de parcourir la MIB à l'aide d'un **browser de MIB**. Cet outil permet de visualiser l'arborescence et d'interroger ou de modifier chaque variable.

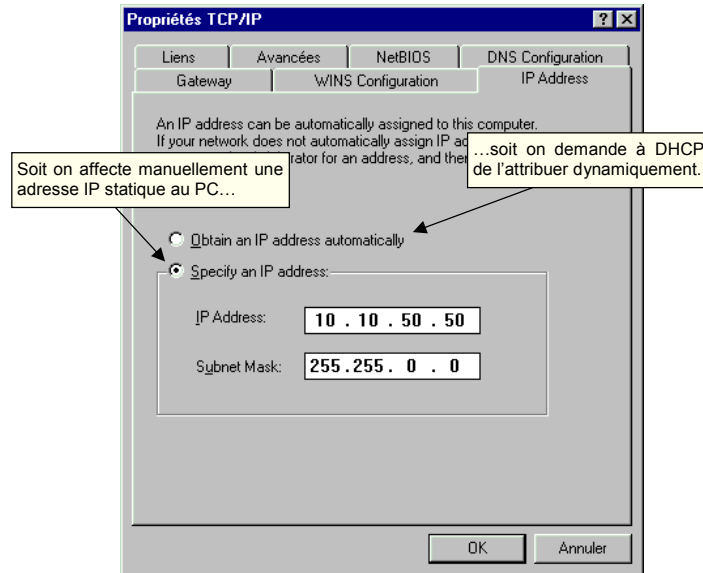


Grâce au browser de MIB, il est possible d'interroger une variable en particulier, puis de programmer des actions automatiques, comme une interrogation périodique du débit entrant et sortant d'une interface, de manière à suivre l'évolution de la charge d'un lien. Avec une série de valeurs, il sera par la suite possible de produire un graphique.



Configurer automatiquement ses PC

La configuration des PC peut s'avérer fastidieuse et être source d'erreurs : adresses dupliquées, masques incorrects, etc. Le protocole **DHCP** (*Dynamic Host Configuration Protocol*) permet d'automatiser ces tâches à l'aide d'un serveur qui héberge les configurations. La plupart des piles IP, dont celle de Microsoft, intègrent un client DHCP. L'unique configuration nécessaire sur les PC consiste à indiquer l'option "Obtenir l'adresse IP par un serveur DHCP" lors de la configuration de la carte réseau.



Par défaut, le PC envoie sa requête à tous les serveurs DHCP et sélectionne généralement le premier qui répond. Il est cependant possible de choisir le serveur DHCP en modifiant la clé de registre "HKEY_local_machine\System\CurrentControlSet\Services\VxD\DHCP\DhcpInfo00\DhcpIPAddress".

Il s'agit de la même configuration que nous avons réalisée au premier chapitre lorsque nous nous sommes connectés à l'Internet. En effet, les fournisseurs d'accès Internet (les ISP) utilisent systématiquement un serveur DHCP pour attribuer les adresses IP aux PC qui se connectent à leur réseau *via* un modem. Il s'agit généralement d'une adresse publique prise dans le plan d'adressage affecté officiellement à l'opérateur par le NIC (*Network Information Center*).

Quelle utilisation de DHCP ?

Que vous disposiez de 10 ou 10 000 postes de travail, DHCP sera tout aussi simple à configurer et, dans tous les cas, il vous facilitera la vie.

Ce protocole permet, avant tout, d'affecter une adresse IP à une station pendant une durée limitée. À chaque initialisation — et lorsque la période de validité est expirée —, le PC demande une nouvelle adresse. Cela procure plusieurs avantages :

- Il n'y a plus de risque d'erreur lié à une configuration manuelle.
- Lorsqu'un PC est déplacé et qu'il change de réseau IP, il n'est plus nécessaire de modifier son adresse IP.
- En considérant que tous les PC ne se connectent pas en même temps, on peut utiliser un pool de 253 adresses (une classe C) pour connecter 500 PC, par exemple. Un ratio de un pour quinze est généralement utilisé par les ISP. Cela permet de pallier la pénurie d'adresses publiques.

On peut également affecter l'adresse de façon permanente, mais on perd alors tous les avantages énumérés précédemment.

Plus intéressant, l'utilisation de DHCP peut être étendue à la configuration de tous les paramètres réseau du PC (liés à la famille TCP/IP ou à d'autres protocoles), tels que le routeur par défaut, le masque IP ou encore le TTL par défaut. Cela procure de nouveaux avantages pour l'administrateur réseau :

- Tous les équipements réseau disposent des mêmes paramètres, ce qui assure une meilleure stabilité de fonctionnement de l'ensemble.
- Tout changement de configuration réseau est automatisé. Des opérations complexes, telles que la migration vers un nouveau plan d'adressage ou l'application d'un paramètre TCP permettant d'optimiser le réseau, sont rendues extrêmement simples et rapides.

Les matériels réseau (routeurs, agents SNMP, etc.) ainsi que les serveurs doivent disposer d'adresses fixes, car ils doivent être connus de tous. Ils peuvent faire appel à DHCP pour obtenir une adresse permanente que vous aurez préalablement réservée ou pour obtenir des paramètres de configuration IP.

LES OPTIONS DHCP (RFC 2132)

La RFC 2132 précise les principales options qui peuvent être affectées par un serveur DHCP. Parmi les plus importantes, on trouve (les numéros d'options sont indiqués entre parenthèses) :

- le masque de l'adresse IP (004) ;
- l'adresse IP des serveur DNS et le nom du domaine DNS dans lequel est situé la station (006 et 015) ;
- le nom de la station ;
- l'adresse IP des serveurs WINS et le type de nœud Netbios (044 et 046) ;
- des paramètres IP, TCP et ARP tels que le MTU (026), le TTL (023), la durée du cache ARP (035), etc. ;
- des routes statiques par défaut ainsi que l'adresse du routeur par défaut (033 et 003) ;
- les serveurs de messagerie SMTP et POP (069 et 070) ;
- divers serveurs par défaut tels que web (072), News (071), NTP (042), etc. ;
- des paramètres relatifs à DHCP (durée de validité de l'adresse, etc.) ;
- les types de messages DHCP (DISCOVER, REQUEST, RELEASE, etc.).

D'autres RFC peuvent décrire des options spécifiques (par exemple la RFC 2244 pour des paramètres Novell). La liste exhaustive des options officiellement reconnues est disponible sur <http://www.iana.org>.

Certains concentrateurs ou commutateurs peuvent télécharger leur système d'exploitation (un fichier exécutable appelé image de boot) ou un fichier de configuration en utilisant le sous-ensemble **BOOTP** également pris en charge par le serveur DHCP.

Avant de commencer, vous pourrez prendre en compte les recommandations suivantes :

- Installez au moins un serveur par site pour des questions de performance et de charge sur les liaisons WAN.
- Pour des questions de sécurité (surtout lorsque le nombre de PC est important), il est préférable de configurer deux serveurs par site, chacun gérant un pool d'adresses.
- La durée de validité des paramètres doit être limitée dans le temps dans le cas où les stations ne sont jamais éteintes (ce qui est le cas des serveurs, par exemple). Une durée de 12 ou 24 heures permet de couvrir une journée de travail et de diffuser de nouveaux paramètres assez rapidement. Pour les serveurs et équipements réseau, une durée plus longue peut être définie, mais l'application d'un nouveau paramètre prendra plus de temps. Vous pourrez de toute façon modifier la durée à tout moment.

Enfin, toutes les piles IP ne prennent pas l'ensemble des options possibles en charge. Il convient donc de vérifier que celle que vous utilisez accepte les options que vous voulez distribuer *via* DHCP.

Comment configurer un serveur DHCP ?

Pour installer le serveur sous Windows NT, il faut se rendre dans la configuration des services IP : “ Démarrer → Paramètres → Panneau de Configuration → Réseau → Services → Ajouter ”.

Pour configurer le serveur DHCP, cliquez sur “ Démarrer → Programmes → Outils d'administration → Gestionnaire DHCP ”. Nous prenons l'exemple de Windows NT, mais le principe est le même sous Unix.

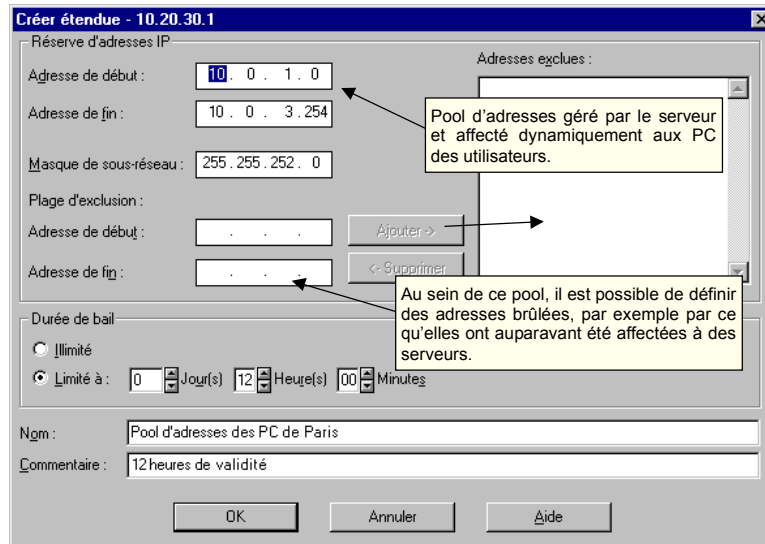
Définir les pools d'adresses

La première étape consiste à définir des pools d'adresses dans lesquels le serveur va piocher pour les affecter aux stations qui en feront la demande. Conformément à notre plan d'adressage, nous avons découpé notre espace d'adressage en trois parties.

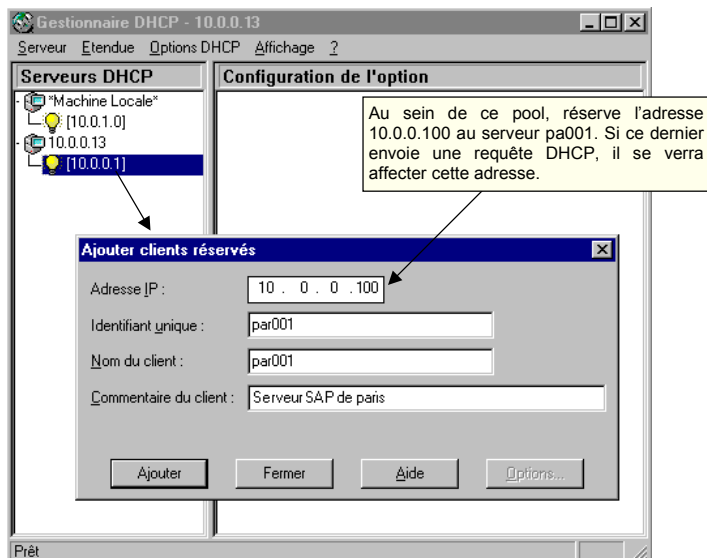
Plage d'adresses	Affectation
De 0.1 à 0.64	Équipements réseau (routeurs, hubs, switches, etc.). Les adresses seront fixes, et seules les options seront distribuées (éventuellement des images de boot). Cela implique de référencer tous les équipements concernés (noms et adresses MAC notamment). Durée de validité des options = 1 semaine.
De 0.65 à 0.255	Serveurs NT, Unix, etc. Les adresses seront fixes, et seules les options seront distribuées. Les options peuvent être communes à tous les serveurs. Durée de validité des options = 1 semaine.
De 1.0 à 3.254	Postes de travail (PC, etc.). Adresses et options affectées dynamiquement. Durée de validité = 12 heures.

Rappelons qu'il est souvent souhaitable de limiter le découpage à deux tranches, une pour les équipements réseau et serveur, et une pour les postes de travail. Les deux premières parties peuvent donc être fusionnées.

À Paris, le serveur DHCP prendra ainsi en charge la plage d'adresses allant de 10.0.0.1 à 10.0.3.254, découpée en deux pools (appelés *scope* ou *étendue* chez Microsoft). Le *scope* dédié aux stations commencera à 10.0.1.0. En cas de saturation de la première tranche, il sera toujours possible de modifier la plage d'adresses affectée à ce *scope*.



La configuration du pool d'adresses pour les équipements réseau et les serveurs nécessite, en plus, de réserver les adresses *via* le menu "Etendue→Ajouter adresse réservée". Avec les clients Windows, le seul moyen d'identifier les stations est d'utiliser l'adresse MAC de la carte réseau. La norme prévoit cependant que l'identifiant puisse être une chaîne de caractères quelconque, le nom de l'utilisateur ou du PC, par exemple.



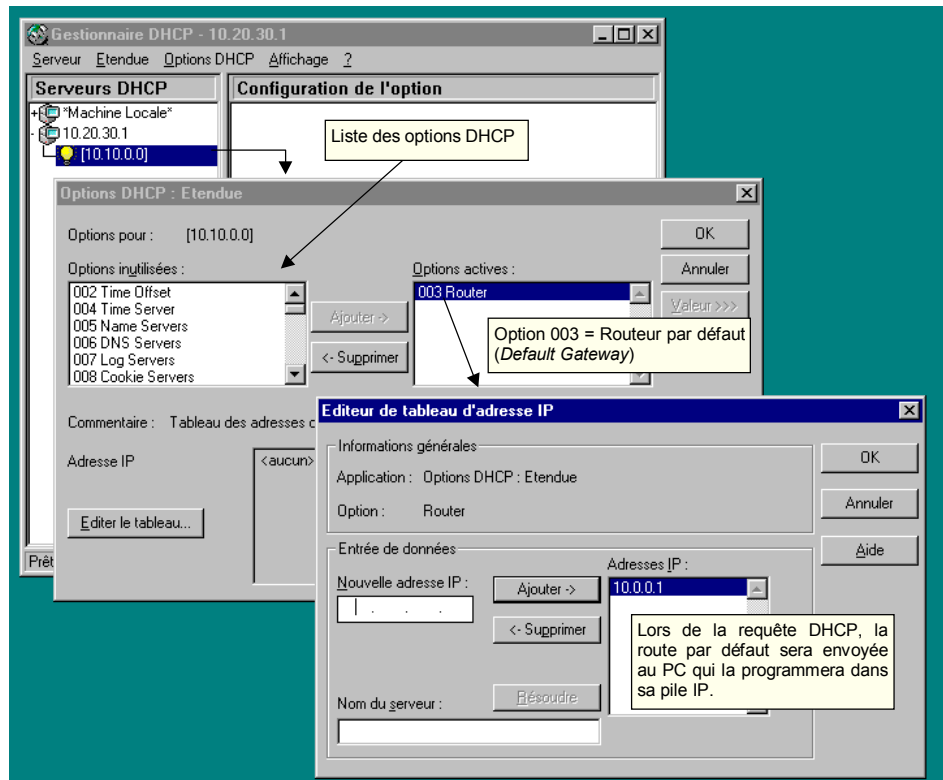
Cette opération peut être fastidieuse, car il faut relever les adresses MAC de tous les équipements concernés (qui sont cependant en nombre moins élevé que les PC). De plus, lorsqu'une carte réseau est changée, il faut mettre à jour la base de données DHCP (c'est une opération en principe peu fréquente, mais il faut y penser le jour où cela arrive).

Si vous ne voulez pas utiliser DHCP pour cette classe d'équipements, il suffit de ne pas ajouter de pool d'adresse. La plan d'adressage facilite votre choix.

Définir les options à distribuer

Les options peuvent être définies à trois niveaux : soit globalement pour tous les pools d'adresses, soit pour chaque pool (*scope*), soit encore individuellement pour chaque client.

Les options communes à tous les nœuds du réseau — par exemple le TTL par défaut ou l'adresse d'un serveur NTP servant de référence à la mise à l'heure des horloges — peuvent être définies globalement dans le menu “ Option DHCP → Global ”.





Si vous voulez affecter comme passerelle par défaut l'adresse IP de la station elle-même (voir chapitre 8), il faut ajouter et positionner à "1" la clé de registre suivante au niveau du pool "HKEY_local_machine\System\ CurrentControlSet\ Services\ DHCPserver\ Subnets\ a.b.c.d\ SwitchedNetworkFlag", où a.b.c.d est l'adresse IP du pool.

LE POINT SUR DHCP (RFC 2131)

DHCP (*Dynamic Host Configuration Protocol*) permet à une station d'obtenir l'intégralité de ses paramètres IP (plus de 65 options recensées à ce jour), ce qui épargne à l'administrateur de devoir configurer manuellement chaque poste de travail. DHCP est une extension du protocole **BOOTP** ; il utilise le même format de paquet.

8 bits		8 bits		8 bits		8 bits	
1 = Requête 2 = Réponse		Type d'adresse physique (1=MAC Eth)		Longueur de l'adresse physique		Saut incrémenté de 1 par les routeurs	
XID : numéro identifiant de manière unique la transaction. La réponse du serveur doit contenir le même XID que la demande du client							
Nombre de secondes depuis que le client a initialisé sa demande			B	Indicateurs (non utilisés) B = bit de Broadcast			
ciaddr - Adresse IP du client si celui-ci la connaît Il peut la connaître s'il demande une prolongation de l'affectation de l'adresse.							
yiaddr - Adresse IP affectée par le serveur							
siaddr - Dans son message DHCP_OFFERT, le serveur DHCP indique ici son adresse IP que le client devra indiquer en retour dans son message DHCP_REQUEST							
giaddr - Adresse IP du routeur ayant relayé le message DHCP. Si cette adresse est non nulle, le serveur sait que la requête a traversé au moins un routeur.							
chaddr - Adresse physique du client (16 octets) Dans le cas d'Ethernet, il s'agit de l'adresse MAC							
sname - Nom optionnel du serveur (64 octets maximum terminés par un 0)							
file - Nom optionnel du fichier à télécharger (128 octets maximum terminés par un 0)							
Options - Liste des paramètres supplémentaires affectés au client (les options sont listées dans le RFC 2132) ainsi que le type de message DHCP (DISCOVER, REQUEST, etc.).							

Si la pile IP (et notamment le module ARP) peut fonctionner sans adresse IP, le **bit de broadcast** peut être mis à 0 dans les requêtes DHCP, ce qui permet au serveur de renvoyer ses réponses dans des trames unicast (à l'adresse MAC indiquée par le client dans le champ *chaddr*). Dans le cas contraire, le bit de broadcast est positionné à 1 par le client, et le serveur répond dans des trames de broadcast MAC (FF:FF:FF:FF:FF:FF).

Cependant, si le champ *giaddr* est non nul, cela veut dire que la requête a transité par un routeur. Le serveur envoie alors le paquet DHCP à cette adresse IP (et donc à l'adresse MAC du routeur *via* ARP). Le port UDP de destination est alors 67 (celui du serveur) — au lieu de 68 qui désigne le client —, ce qui permet au routeur d'identifier les paquets DHCP à traiter (voir plus loin).

Si le client possède déjà une adresse IP (champ *ciaddr* non nul), il peut demander des paramètres de configuration complémentaires (les **options DHCP**) en envoyant le message DHCP_INFORM. Le serveur envoie alors sa réponse à l'adresse IP indiquée (donc dans une trame MAC unicast).

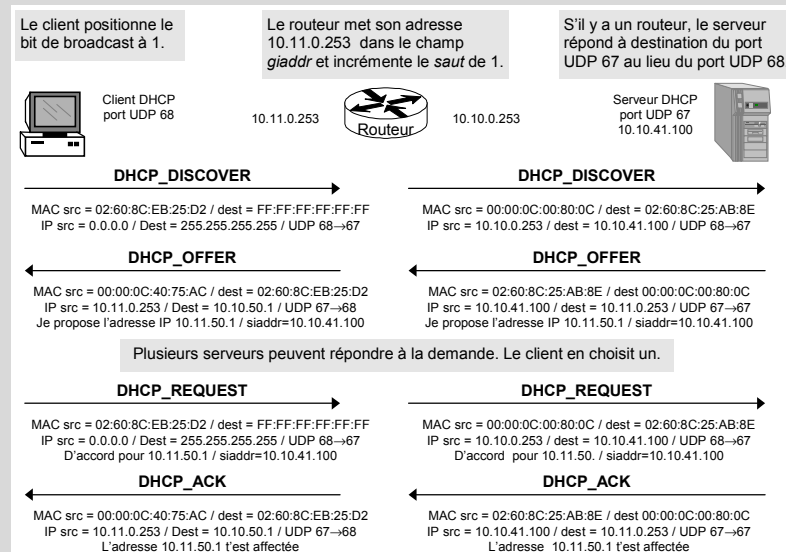
...

LE POINT SUR DHCP (SUITE)

Un client effectue sa requête en deux temps :

- Tout d'abord, il recherche un serveur DHCP, et attend les offres du ou des serveurs.
- Ensuite, il confirme sa demande auprès du serveur qu'il a choisi, et attend une réponse lui confirmant que l'adresse IP a bien été réservée.

Les options sont également négociées au cours de cet échange : le client indique celles déjà configurées dans sa pile IP, et les serveurs lui proposent les leurs.



Le client envoie toujours ses requêtes dans des trames de **broadcast MAC** pour plusieurs raisons :

- La requête initiale (DHCP_DISCOVER) permet de découvrir plusieurs serveurs (un serveur principal et un serveur de secours).
- Lors de la requête de confirmation (DHCP_REQUEST), la plupart des piles IP ne peuvent activer le module ARP sans adresse IP.
- Le client ne sait pas si le serveur est situé sur le même réseau ou s'il est séparé par un routeur. Dans ce dernier cas, si la trame est destinée à l'adresse MAC du serveur, elle ne traversera pas le routeur, sauf s'il fonctionne en mode proxy ARP (voir chapitre 8), ce que la station ne peut présupposer.

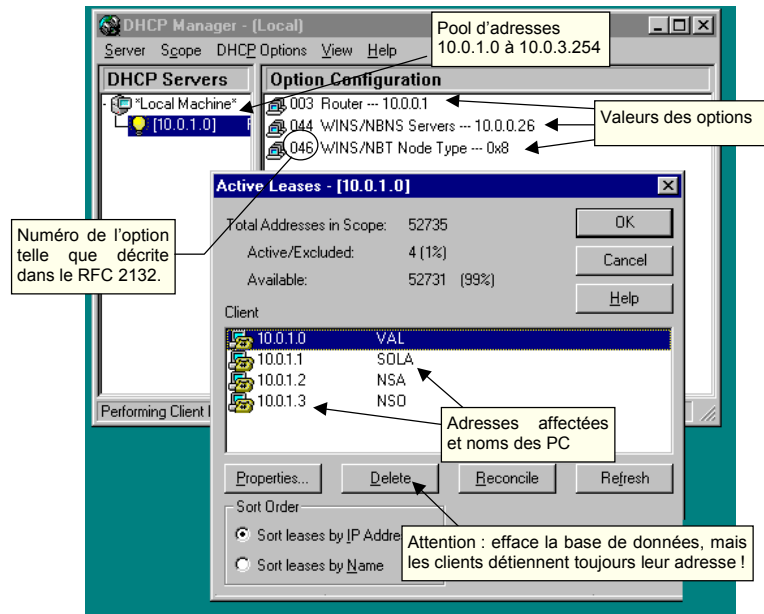
Configuré en relais DHCP/BOOTP, le routeur convertit les broadcast MAC/IP en adresses unicast à destination du serveur DHCP ou BOOTP.

À l'approche de l'expiration de la période de validité, le client demande à renouveler son bail auprès du serveur (DHCP_REQUEST) en indiquant son adresse IP dans le champ *ciaddr*. Le serveur peut alors proposer la même adresse, une nouvelle, ou encore accepter celle demandée par le client.

En principe, c'est au client qu'il incombe de vérifier que l'adresse allouée par le serveur n'est pas utilisée par une autre station (le serveur peut, en effet, être situé de l'autre côté d'un routeur). Le client génère à cet effet une requête ARP sur l'adresse qui vient de lui être allouée.

La même option peut être définie à plusieurs niveaux, mais les options individuelles ont priorité sur les options d'un pool, qui elles-mêmes ont priorité sur les options globales.

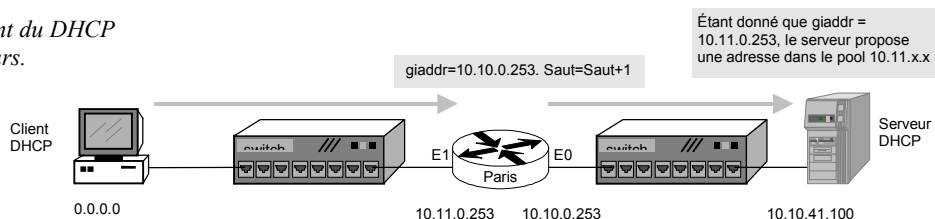
L'exemple suivant montre les options définies pour un pool, ainsi que les adresses déjà affectées à des clients.



Configurer les routeurs

Si les stations sont situées sur un réseau IP autre que le serveur DHCP, les requêtes DHCP doivent transiter par un routeur. Or, ce type d'équipement ne transmet jamais les trames de broadcast MAC, car il est justement conçu pour délimiter les domaines de broadcast (voir l'encart « Le point sur Ethernet » au chapitre 6). Il faut donc configurer explicitement les routeurs afin de pouvoir relayer les requêtes DHCP.

Figure 16-3.
Fonctionnement du DHCP avec les routeurs.



```
interface ethernet 1
ip helper-address 10.10.41.100
ip helper-address adresses ip d'autres serveurs DHCP
```

Avec la commande précédente, les trames de broadcast MAC dont le port UDP est égal à 67 seront transmises dans une trame unicast (*via* la résolution ARP) à destination du serveur DHCP ou BOOTP et, en retour, vers le client.

LE POINT SUR BOOTP (RFC 951 ET 1542)

BOOTP (*Bootstrap Protocol*) permet à un équipement réseau d'obtenir son adresse IP ainsi que le nom d'un fichier à télécharger *via* TFTP (*Trivial File Transfer Protocol*). Il peut s'agir d'un **fichier de configuration** ou d'un **exécutable** (appelé image de boot), tel qu'un système d'exploitation ou un micro-code. Ce protocole ne permet pas d'affecter dynamiquement les adresses, et nécessite donc de connaître les adresses MAC ou d'affecter un nom aux équipements qui émettent des requêtes.

DHCP a repris exactement les mêmes spécifications que BOOTP en étendant ses possibilités. Un serveur DHCP prend en charge les requêtes BOOTP (compatibilité ascendante), alors que l'inverse n'est pas possible.

Des équipements réseau qui ne disposent pas de mémoire flash, tels que des concentrateurs, des commutateurs ou des serveurs d'accès distants, utilisent BOOTP pour télécharger leur code exécutable.

Pour la petite histoire, la RFC 1542, relative aux extensions de BOOTP (l'équivalent des options DHCP), discute de l'utilité du bit de broadcast en rappelant le paradigme de la poule et de l'œuf : une station ne peut pas fonctionner sans adresse IP ; cependant, elle envoie et reçoit des paquets IP qui doivent justement lui permettre d'obtenir cette adresse IP ; mais elle ne peut pas traiter ces paquets puisqu'elle ne dispose pas d'adresse IP, etc.

Installer plusieurs serveurs

Si plusieurs serveurs sont utilisés (en partage de charge et en redondance), le pool d'adresses doit être découpé afin d'éviter les doubles affectations. La répartition peut se faire à parts égales, comme suit :

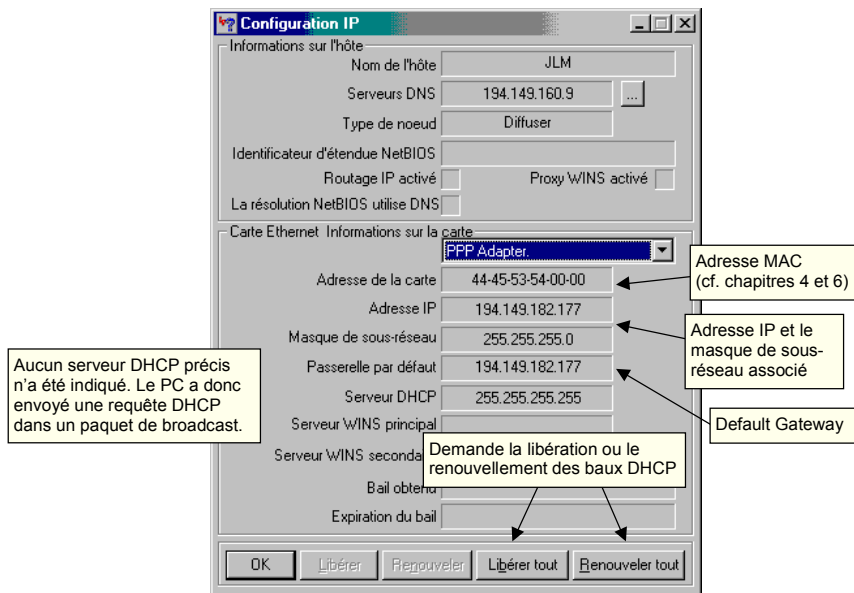
- le premier serveur affecte les adresses comprises entre 10.0.1.0 et 10.0.2.127 ;
- le second serveur affecte les adresses comprises entre 10.0.2.128 et 10.0.3.254.

Les adresses affectées de manière fixe doivent être réservées de manière identique sur chaque serveur.

Le serveur de Microsoft ne permet pas de mettre en place un réel partage de charge avec une redondance complète. Pour cela, d'autres serveurs DHCP plus perfectionnés existent sur le marché.

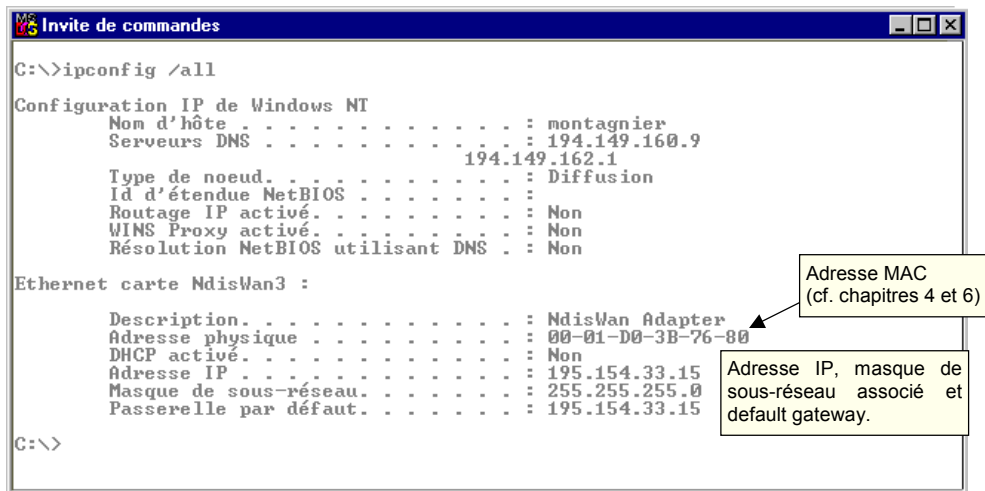
Vérifier la configuration de son PC

Plusieurs utilitaires permettent de vérifier le paramétrage TCP/IP de son PC. Sous Windows 9.x et Me, il s'agit de la commande **winipcfg**.



Sous Windows NT, la commande équivalente est **ipconfig** :

- **ipconfig /all** affiche tous les paramètres réseau.
- **ipconfig /release** envoie un message DHCP_RELEASE au serveur pour libérer l'adresse IP.
- **ipconfig /renew** envoie un DHCP_REQUEST pour demander le prolongement de la validité de l'adresse IP, ou un DHCP_DISCOVER si la station ne possède pas d'adresse.



Sous les deux environnements, la commande **route** permet de visualiser la table de routage de la pile IP.

Route permanente (sauvegardée dans la Registry).

Ajoute une route statique : pour envoyer un paquet vers le réseau 192.168.0.0, le PC le transmettra au routeur 10.0.0.253.

Route par défaut (default gateway)

Adresse de loopback utilisée pour vérifier que la pile IP fonctionne bien (par exemple avec un ping)

Groupe multicast par défaut (cf. chapitre 12)

Adresse réseau	Masque réseau	Adresse passerelle	Interface	Métrique
0.0.0.0	0.0.0.0	10.10.0.1	10.10.50.50	1
127.0.0.0	255.0.0.0	127.0.0.1	127.0.0.1	1
192.168.0.0	255.255.255.0	10.10.0.253	10.10.50.50	1
10.10.0.0	255.255.0.0	10.10.50.50	10.10.50.50	1
10.10.50.50	255.255.255.255	127.0.0.1	127.0.0.1	1
10.10.255.255	255.255.255.255	10.10.0.1	10.10.50.50	1
224.0.0.0	224.0.0.0	10.10.0.1	10.10.50.50	1
255.255.255.255	255.255.255.255	10.10.0.1	10.10.50.50	1

Enfin, la commande **Netstat** permet de vérifier le bon fonctionnement de la couche TCP/IP. Elle permet, par exemple, de visualiser les connexions actives.

Noms obtenus par le DNS et correspondant à l'adresse IP cible.

Numéros de ports TCP source et destination.

Proto	Adresse locale	Adresse extérieure	Etat
TCP	montagnier:1025	localhost:1026	ETABLIE
TCP	montagnier:1026	localhost:1025	ETABLIE
TCP	montagnier:1328	ADSERVER58:80	ETABLIE
TCP	montagnier:1641	cctvw02.ctec.com:80	TEMPS D'ATTENTE
TCP	montagnier:1666	cio-sys.cisco.com:80	ETABLIE
TCP	montagnier:1667	cio-sys.cisco.com:80	ETABLIE
TCP	montagnier:1676	www.smartbooks.com:80	ATTENTE_FERMER
TCP	montagnier:1677	www.smartbooks.com:80	ATTENTE_FERMER
TCP	montagnier:1678	www.smartbooks.com:80	ATTENTE_FERMER
TCP	montagnier:1679	www.smartbooks.com:80	ATTENTE_FERMER

17

La gestion des noms

Partie de la câblerie et de la basse filasse, nous abordons dans ce chapitre les couches applicatives, c'est-à-dire les services réseau. Pour ainsi dire, nous nous élevons dans les couches supérieures du réseau.

Car, de nos jours, l'administrateur réseau ne peut pas se contenter de fournir le transport des données. Il doit en faciliter l'accès à ses utilisateurs.

En outre, plus le nombre d'utilisateurs est élevé, plus il est important de simplifier les tâches administratives.

Il convient donc d'utiliser des outils qui simplifient la vie des utilisateurs et celle des exploitants. Le service de nom, ou DNS (*Domain Name System*), est le premier d'entre eux.

Dans ce chapitre vous apprendrez ainsi :

- à comprendre le fonctionnement du DNS ;
- à définir un plan de nommage ;
- à configurer les serveurs DNS ;
- à configurer les PC ;
- à interroger la base de données du DNS.

Présentation du DNS

Pour vos utilisateurs et vous-même, il serait bien plus pratique d'utiliser des noms de machines plutôt que des adresses, à l'instar de ce qui se fait sur l'Internet. De même que l'on accède à `www.3com.com`, il serait bien plus simple d'accéder à votre serveur au moyen du nom `www.societe.fr` plutôt que de son adresse IP.

La première solution repose sur l'utilisation des fichiers `hosts` (localisés dans `/etc` sous Unix, et dans `\Windows` sous Windows 9x). Ce fichier contient simplement la correspondance entre adresses IP et noms de machines :

```
10.0.0.1      par001
10.0.0.100    nt001
10.0.0.101    mail
```

L'inconvénient de cette méthode est qu'il faut configurer le fichier sur chaque poste de travail, et ce, à chaque changement d'adresse ou de nom. On pourrait imaginer une distribution automatique de ce fichier, mais cette opération serait très complexe et source d'erreur avec les PC. De plus, l'espace de nommage est "plat" (tous les noms sont au même niveau).

La seconde solution, de loin la meilleure, repose donc sur l'utilisation d'un système DNS (*Domain Name System*). C'est celle utilisée à grande échelle sur l'Internet (plusieurs millions de machines référencées début 2001) et qui convient également pour vos 10 ou 10 000 PC.

Les composants du DNS

Le DNS a déjà été introduit aux chapitres 3, 15 et 16 lorsque, par exemple, nous nous sommes promenés sur l'Internet. Il s'agit ici de recréer un DNS privé, c'est-à-dire réservé à nos utilisateurs.

Il faut pour cela définir :

- un espace de nommage hiérarchique découpé en domaines ;
- des serveurs gérant des bases de données ;
- des clients appelés *resolver* ;
- un protocole d'échange entre clients et serveurs d'une part, et entre serveurs d'autre part.

Tous ces composants sont décrits dans une série de RFC dont les premiers remontent à 1987.

Élaborer un plan de nommage

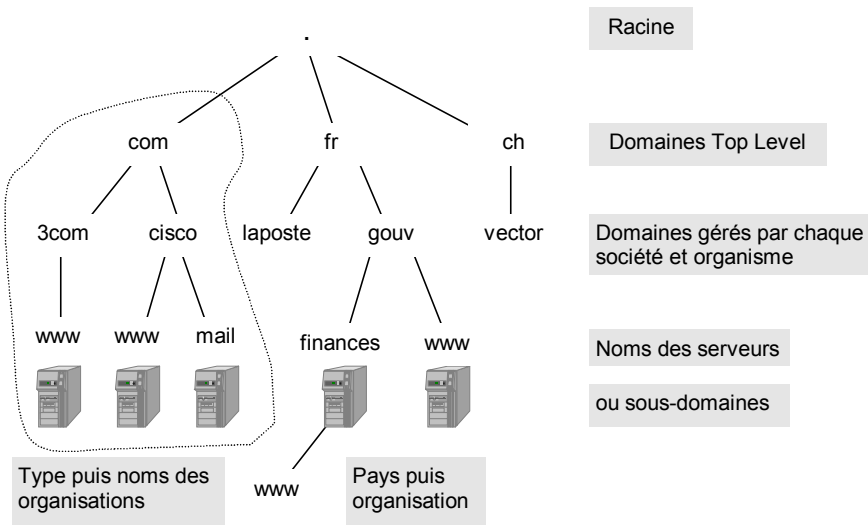
Le nommage DNS est organisé sous forme d'arbre, avec une racine et des domaines qui lui sont rattachés. Le plan de nommage consiste donc à définir cette arborescence et la manière d'affecter des noms aux objets (les feuilles de l'arbre).

Définir l'arborescence DNS

Dès le début de votre réflexion, vous serez confronté au dilemme classique : définir une arborescence qui reflète l'organisation de la société ou une arborescence qui reflète son implémentation géographique ?

Par expérience, l'une n'est pas meilleure que l'autre, car toutes deux sont soumises aux aléas des changements. L'approche organisationnelle est soumise au changement du nom de la société (suite à une décision stratégique, à un rachat, etc.) ou du service (suite à une réorganisation), tandis que l'approche géographique est soumise aux déménagements.

L'Internet a d'ailleurs retenu les deux approches.



Pour notre DNS privé, donc à usage purement interne, nous n'avons pas besoin de faire référence au pays et au nom de la société, mais plutôt à la ville, au nom du site et au nom des directions (ou, selon la terminologie propre à chaque société, des divisions, des *Business Units*, etc.).

Nous pouvons cependant retenir la même approche mixte en fonction du degré de centralisation de chaque département.

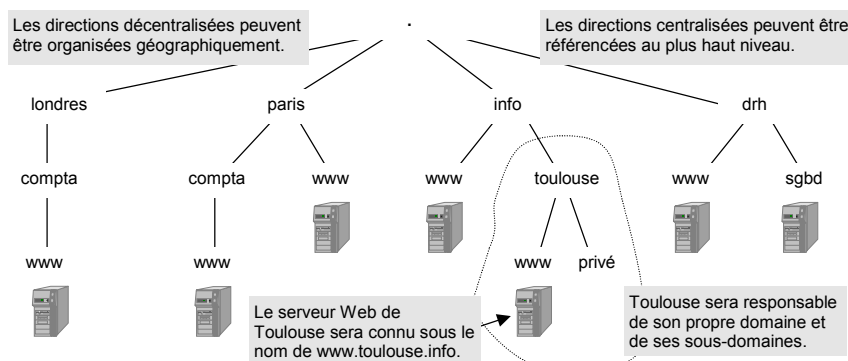
Représentation dans le DNS	Organisation centralisée	Organisation décentralisée
Racine	Par convention, le point ('.')	Par convention, le point ('.')
Domaine Top Level	Nom de la direction	Nom de la ville
Domaine (optionnel)	Nom de la ville	Nom de la direction

Il est conseillé de ne retenir que les invariants ou, pour être exact, les éléments qui sont susceptibles de changer le moins souvent. Dans notre cas, ce sont les directions (c'est-à-dire les services situés au sommet de la hiérarchie organisationnelle) et les villes principales où est implantée notre société.

L'espace de nommage est indépendant de la localisation géographique : tous les immeubles d'une même ville peuvent donc être référencés dans le même domaine DNS.

Par exemple, l'indication d'éléments pouvant changer fréquemment, tels que le nom du service au sein d'une direction ou le nom du site dans une ville, est source de complication, et entraîne un travail supplémentaire de reconfiguration. Plus vous ajouterez de niveaux au DNS, plus vous devrez effectuer de mises à jour. En revanche, si un service au sein d'une direction devait garder son autonomie, on pourrait envisager de lui déléguer la gestion de son sous-domaine.

Figure 17-1.
*Définition
d'un DNS privé.*



Le DNS offre donc beaucoup de souplesse :

- Une même machine (ayant une adresse IP) peut être référencée dans plusieurs domaines.
- Une même machine peut être référencée sous plusieurs noms principaux et/ou sous un nom principal associé à des alias.
- La gestion d'un domaine peut être déléguée à un nouveau service devenu autonome, ou être reprise de façon centrale.
- Une direction peut gérer son propre DNS (avec sa propre racine).
- Les machines peuvent être référencées dans plusieurs DNS.
- Un DNS peut comporter une centaine de niveaux.

On le voit, le DNS permet toutes sortes de fantaisies. Il est donc de votre responsabilité d'en assurer la cohérence et la simplicité. Ainsi, les fonctionnalités présentées ci-dessus doivent-elles plutôt être utilisées pour faciliter les périodes de transition lors de changements d'organisation ou de déménagement ou, d'une manière générale, pour répondre à des situations exceptionnelles.

Standardiser le nommage des objets

Il n'est pas nécessaire de mettre les postes de travail dans la base DNS, car ces derniers ne sont que des clients ; ils ne sont pas connus en tant que serveur. Seuls sont renseignés dans les bases DNS les équipements réseau (pour les exploitants) et les serveurs (pour les utilisateurs).

Il est conseillé d'adopter un codage différent et adapté à chaque type d'objet, c'est-à-dire à sa nature et à sa fonction. Par exemple, les serveurs ne seront pas nommés de la même manière que les routeurs. Le nom principal doit correspondre à un invariant, et les alias à la particularité du moment.

Si la sécurité est privilégiée et que le nom ne doit pas permettre d'identifier ces éléments, des noms neutres peuvent être retenus (noms de musiciens, de fleurs, de planètes, etc.).

Objet	Invariant	Variant
Serveur	Système d'exploitation (Unix, NT)	Fonction (web, messagerie, base de donnée, etc.)
Routeur	Marque (Cisco, 3com, etc.) et fonction de routage	Localisation
Concentrateur	Marque (3com, etc.) et fonction de concentration	Localisation et, éventuellement, emploi de réseaux différents (Ethernet, ATM, etc.)
Commutateur	Marque (Cisco, 3com, etc.) et fonction de commutation	Localisation et, éventuellement, emploi de réseaux différents (Ethernet, ATM, etc.)
Serveur d'accès distant	Marque (Cisco, 3com, etc.) et fonction d'accès distant	Localisation

La localisation peut être considérée comme étant un variant de faible impact à partir du moment où l'équipement doit de toute façon être reconfiguré (changement d'adresse IP, par exemple).

Il faut privilégier un nommage simple des équipements auxquels on accède le plus fréquemment : les serveurs (auxquels accèdent les utilisateurs) et les routeurs et serveurs d'accès distants (auxquels accèdent les exploitants réseau). Le nom doit :

- Ne comporter que des caractères alphanumériques (minuscules et/ou majuscules) et des tirets (-). Ils constituent, en effet, le plus petit dénominateur commun dans le monde de l'informatique. Les autres caractères sont à proscrire, car certains systèmes ne les acceptent pas.
- Être court, afin d'être facile à mémoriser et rapide à saisir au clavier.
- Contenir une ou deux alternances de noms et de chiffres pour en améliorer la lisibilité.
- Contenir un tiret au maximum pour séparer deux champs alphabétiques ou numériques, afin d'en améliorer la lisibilité.

Équipement	Codage retenu	Nom principal (nom système)	Exemple
Serveur NT	Nom neutre (planètes)	mars pluton	mars.marseille. mars.paris.
Unix	Nom neutre (musiciens)	mozart schubert	mozart.toulouse.info.
Routeur	Ville sur trois lettres et numéro d'ordre sur trois chiffres	par001, par002 toul001	par001.info. par001.paris.
Serveurs d'accès distants	Préfixe svc, suivi du codage identique aux routeurs	ras-par001 ras-mar002	ras-mar002.info.

Un serveur peut être connu sous plusieurs noms :

- Un premier qui désigne sa nature et qui correspond au nom système défini lors de l'installation. Il ne change pas (sauf lors d'une réinstallation).
- Un ou plusieurs autres qui désignent sa ou ses fonctions et qui correspondent à une ou plusieurs applications (base de donnée, web, etc.)

Pour les utilisateurs, le meilleur moyen d'identifier et de mémoriser le nom d'un serveur est de lui donner le nom de l'application dont ils se servent. Un **alias** correspondant à la fonction du serveur pourra donc être défini et correspondre à la fonction du serveur (il varie au cours du temps alors que le nom système ne change que si l'on réinstalle complètement la machine).

Fonction	Codage retenu	Alias	Exemple
Serveur web principal	Convention universelle	www	www.paris. www.toulouse.info. www.drh.
Autres serveurs web	Convention suivie d'un numéro d'ordre sur un chiffre	www1 www2	www.paris. www1.paris. www2.paris.
Messagerie	Mail suivi d'un numéro d'ordre	mail mail1	mail.paris. mail1.info.
Autres applications	Nom de l'application	compta rivage	compta.paris. rivage.info.

Certains serveurs centraux, c'est-à-dire communs à l'ensemble de la société, pourront être situés juste sous la racine. On aura, par exemple, www pour le serveur servant de point d'entrée à toute la société : il contiendra les informations du jour et des liens vers les autres serveurs web. On pourra aussi y placer les serveurs qui réalisent l'interconnexion des messageries.

Les concentrateurs et les commutateurs ne sont généralement pas accessibles au moyen d'une connexion Telnet, mais *via* SNMP et une station d'administration. Le codage des noms peut donc être plus complexe. Il peut refléter leur localisation géographique et, éventuellement, leur fonction, afin de faciliter leur identification. Dans notre cas, nous avons choisi le codage suivant : [type]-[ville][étage][immeuble].

Champ	Signification	Code
[type]	Type de matériels, sur une lettre	H = Hub, S = switch, A = ATM, F = Frame Relay, I = FDDI
[ville]	Abréviation de la ville, sur trois lettres	par = Paris, tou = Toulouse tur = Tour, etc.
[immeuble]	Lettre majuscule identifiant un immeuble dans la ville	Le code dépend de la ville D = Descartes, M = Montparnasse, etc.

Ce qui donne, par exemple, h-par05a, s-tou05D, etc.

Afin de faciliter leur lecture et leur traitement dans des bases de données, il est préférable que chaque champ ait une longueur fixe.

Configurer les serveurs DNS

Le DNS est géré par des serveurs qui assurent plusieurs fonctions :

- la gestion d'une **zone**, c'est-à-dire d'un domaine et de ses sous-domaines : on dit que le serveur a **autorité** sur la zone ;
- l'échange des bases de données au sein d'une zone : un serveur est désigné **primaire** pour une zone et distribue la base de données aux serveurs **secondaires** ;
- le **relais** des requêtes DNS d'un client sur un nom situé dans une autre zone ;
- le **cache** des requêtes clients de manière à limiter les requêtes ;
- enfin, un serveur doit être désigné **racine** de l'arbre DNS.

Tout serveur DNS est serveur cache et peut être à la fois ou seulement :

- **primaire** pour un ou plusieurs domaines ;
- **secondaire** pour un ou plusieurs autres domaines ;
- **racine**.

Comme pour le nommage, le DNS offre de nombreuses possibilités et peu de contraintes :

- Il faut au moins un serveur racine.
- Il faut un et un seul serveur primaire par domaine.
- Toutes les modifications de la base de données d'une zone doivent être réalisées sur le serveur qui a autorité (le serveur primaire).

- Il peut y avoir aucun ou autant de serveurs secondaires par domaine.
- Les fonctions primaires et secondaires sont exclusives pour un domaine donné.
- Mais un serveur peut être à la fois primaire pour un domaine et secondaire pour un autre.
- Un serveur DNS peut ne servir que de cache (c'est-à-dire n'être ni racine, ni primaire, ni secondaire).
- L'administrateur peut déléguer la gestion d'un domaine faisant partie de la zone d'autorité à un autre serveur qui devient alors primaire pour le domaine.
- Les clients et les serveurs DNS peuvent être situés n'importe où sur le réseau : sur des lieux géographiques différents et sur des réseaux IP différents.
- Il peut y avoir plusieurs DNS distincts.
- Les postes de travail et les serveurs peuvent être référencés dans n'importe quel domaine.

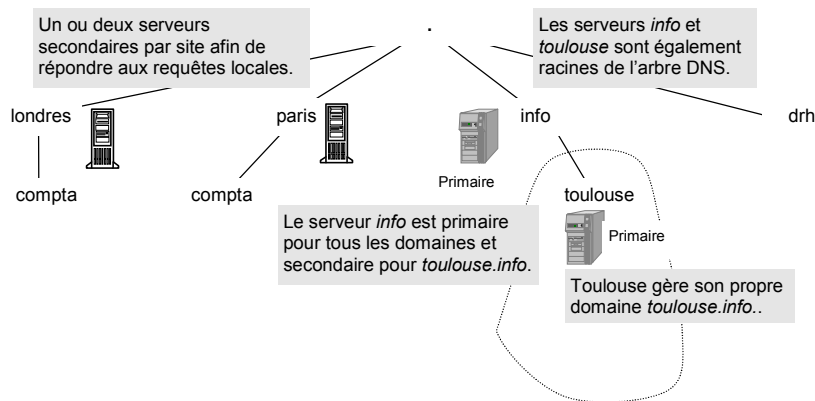
Le DNS permet de faire tout et n'importe quoi. Il convient donc de respecter quelques règles afin de conserver une certaine cohérence au sein de votre société. Ainsi, vous devez prévoir :

- au moins un serveur DNS par site afin de ne pas surcharger les liaisons WAN ;
- au moins un serveur secondaire en partage de charge et en secours ;
- un arbre DNS pour toute la société, même si chacun gère son propre domaine.

Les bonnes pratiques du DNS sont exposées dans la RFC 1912.

Selon ces principes, nous avons décidé de créer l'architecture présentée ci-après.

Figure 17-2.
Les serveurs DNS.

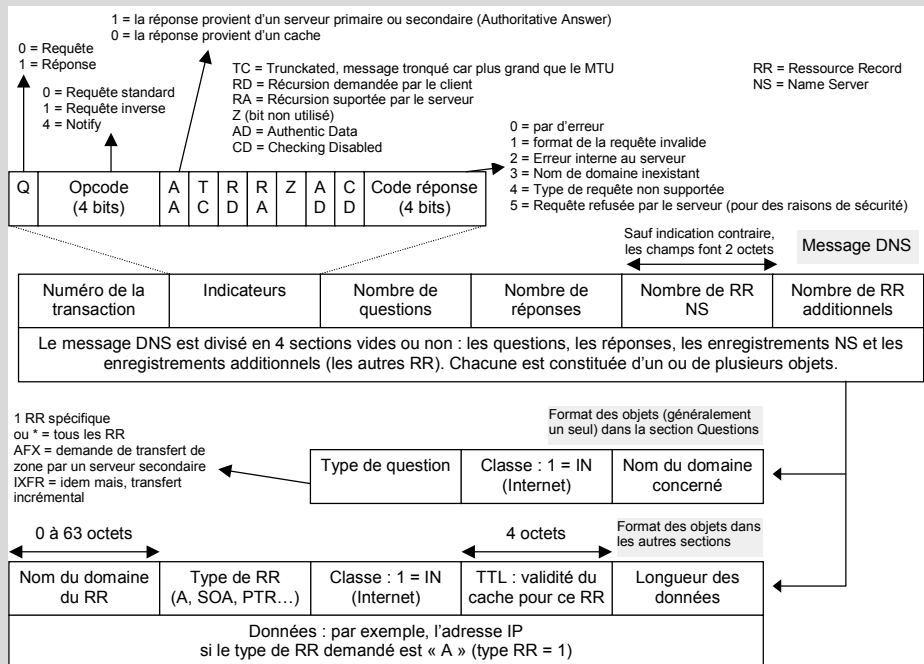


Chaque serveur DNS gère une partie de la base de données DNS, celle représentant la zone sur laquelle il a autorité. Le serveur de Toulouse ne contiendra ainsi que les objets situés dans la zone « *toulouse.info* ».

Sous Windows NT, la manipulation de la base de données est réalisée à l'aide du Gestionnaire DNS, situé dans le menu "Démarrer → Programmes → Outils d'administration". Celui-ci sauvegarde les données dans la base des registres Windows (clé \HKEY_LOCAL_MACHINE\System\CurrentControlSet\Services\Dns) et maintient une copie sous forme de fichiers ASCII dans le répertoire \Winnt\System32\Dns. Ces fichiers respectent le format des serveurs DNS sous Unix, tels que ceux produits par la version appelée BIND et développée à Berkeley.

LE POINT SUR LE DNS (RFC 1034, 1035, 1995, 1996, 2181)

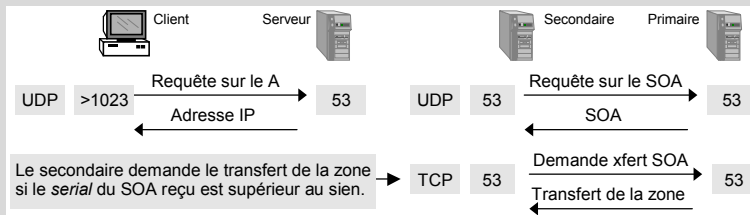
Le DNS (*Domain Name System*, quelquefois appelé *Domain Name Service*) définit une base de données découpée en **zones** (constituées d'un **domaine** et de ses sous-domaines) correspondant à une partie d'un **arbre hiérarchique**. Un serveur peut avoir autorité sur une ou plusieurs zones. S'il est **primaire**, il est maître de la zone. S'il est **secondaire**, il dispose d'une copie qu'il demande régulièrement au [serveur] primaire. S'il est **racine**, le serveur a autorité sur la racine de l'arbre. Les serveurs racines peuvent **déléguer leur autorité** aux serveurs gérant le premier niveau de domaines, appelés domaines **Top Level**, et ainsi de suite. Tous les serveurs font également office de **cache** pour les requêtes DNS émises par les clients appelés **resolver**.



Afin de diminuer la taille des messages DNS, les noms apparaissant plusieurs fois peuvent être remplacés par des pointeurs de deux octets indiquant la position de l'unique exemplaire du nom. La longueur maximale d'un nom d'objet est de 63 octets, et celle d'un nom de domaine complet (y compris celui de l'objet) est de 255 octets.

LE POINT SUR LE DNS (SUITE)

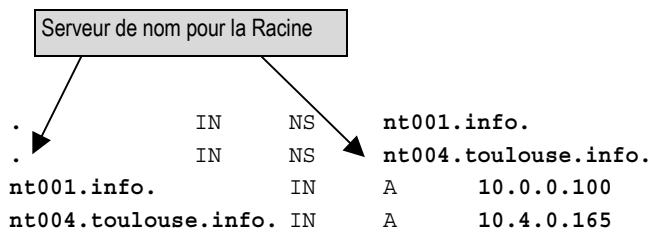
Les messages DNS transitent dans des paquets UDP ou TCP selon les cas (port 53).



Le RFC 1995 précise un mode de transfert incrémental (seules les modifications sont transférées). Le RFC 1996 spécifie, quant à lui, un mécanisme permettant au primaire de notifier au secondaire que le SOA vient d'être modifié. Cela évite d'attendre la fin de la période indiquée dans le paramètre *refresh*.

Configurer le fichier cache

Tous les serveurs participant à votre DNS doivent connaître les serveurs racines. Ces informations résident dans le **fichier cache** (attention, celui-ci n'a rien à voir avec le serveur appelé cache). Sous Windows NT, ce fichier est situé dans le répertoire `\Winnt\System32\Dns\Cache.dns`. Il doit être édité manuellement sur chaque serveur DNS de votre société :



Les instructions NS (*Name Server*) et A (*Address*) indiquent les adresses IP des serveurs racines. Il peut y en avoir autant que nécessaire, afin d'assurer la redondance et le partage de charge.

Si vous voulez construire un DNS intranet directement rattaché à l'Internet, il faut y indiquer les serveurs racines officiels. (Vous pouvez vous procurer la dernière version sur le site <ftp://ftp.rs.internic.net/domain/named.root>.) Vous obtenez alors le fichier cache suivant (extrait) :

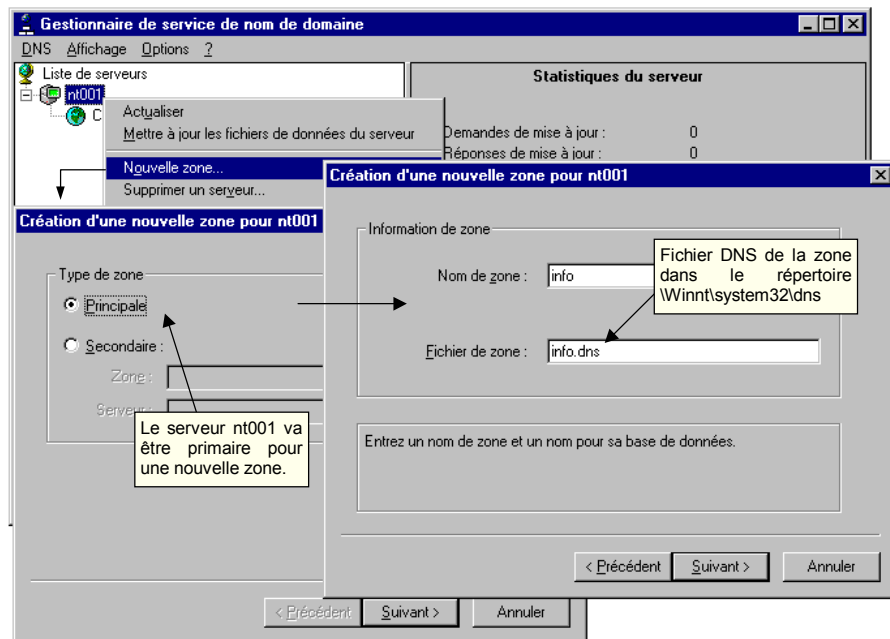
```

;
; formerly NS.INTERNIC.NET
;
.                3600000    IN    NS      A.ROOT-SERVERS.NET.
A.ROOT-SERVERS.NET. 3600000    A      198.41.0.4
;
; formerly NS1.ISI.EDU
;
.                3600000    NS      B.ROOT-SERVERS.NET.
B.ROOT-SERVERS.NET. 3600000    A      128.9.0.107
.....

```

Configurer un serveur primaire

Dans le DNS, il est nécessaire de configurer un serveur primaire unique pour chaque zone. Un serveur primaire pour une zone est primaire pour le domaine situé en tête de cette zone, ainsi que pour tous les sous-domaines.



Cela a pour effet de produire le fichier C:\Winnt\System32\Dns\info.dns :

```

;
; Database file info.dns for info zone.
;   Zone version: 1
;
@           IN      SOA  nt001.info. administrateur.info. (
                                1           ; serial number
                                3600        ; refresh
                                600         ; retry
                                86400       ; expire
                                3600       ) ; minimum TTL

@           NS       nt001
nt001       A        10.0.0.100
nt004.toulouse A      10.4.0.165

```

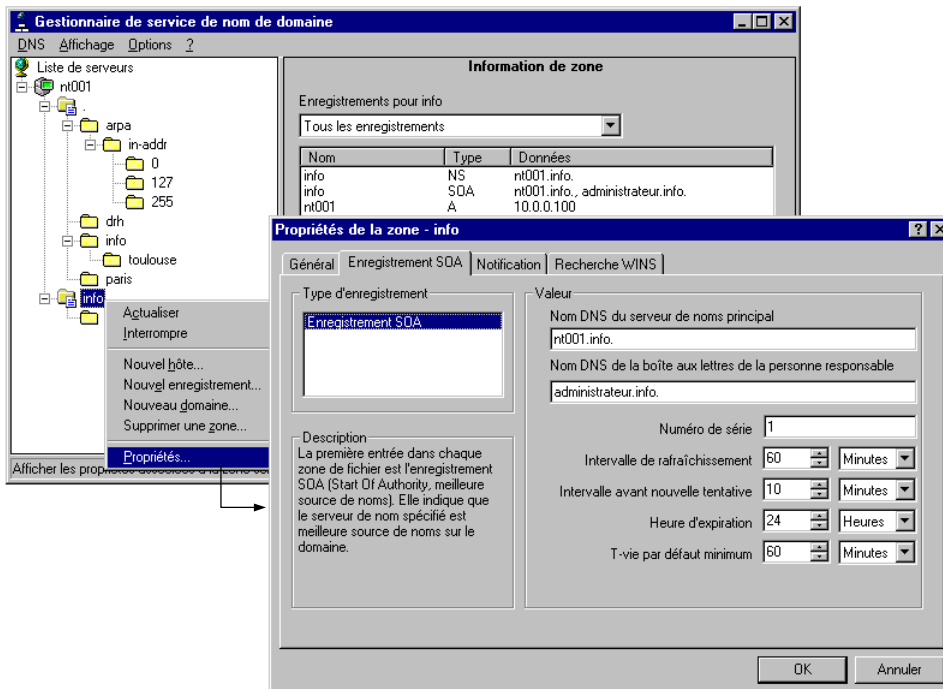
Les parenthèses doivent être utilisées si les paramètres tiennent sur plusieurs lignes.

Description de la zone info.

Référence relative à la zone en cours, ici "info. ".

Serveur primaire pour la zone en cours.

Que ce soit sous NT ou sous Unix, il est possible d'éditer manuellement ce fichier pour modifier les paramètres de la zone. Avec le gestionnaire DNS de Windows, on obtient l'écran suivant.



Les paramètres de la zone (SOA, *Start Of Authority*) sont définis au niveau du serveur primaire pour être ensuite utilisés par les serveurs secondaires. Le champ “ Temps par défaut minimum ” pourra également être utilisé par les resolvers pour déterminer la durée de conservation de l’enregistrement dans leur cache.

Si vous voulez limiter les échanges entre serveurs, mieux vaut positionner le paramètre “ refresh ” à une valeur assez élevée, généralement 12 à 24 heures, sauf durant les périodes de changement. En prévision de telles périodes, vous pouvez réduire cette valeur à 1 heure (voire moins selon vos besoins).

Vous pouvez également activer un mécanisme de notification automatique qui permet au serveur secondaire d’être averti immédiatement de tout changement. À réception de ce signal, le serveur secondaire déclenchera alors un transfert de zone. Il suffit pour cela d’ajouter, dans l’onglet “ Notification ”, les adresses IP des serveurs secondaires.

LA BASE DE DONNÉES DU DNS

La base de données consiste en des fichiers ASCII contenant des **enregistrements** appelés **RR** (*Resource Record*) dont le format générique est le suivant :

nom	[ttl]	[IN]	RR	paramètres
nom	Désigne le nom du domaine. S’il est omis, c’est le même que celui du SOA.			
ttl	Time To Live. Indique pendant combien de secondes le <i>resolver</i> , qui aura fait une requête, pourra conserver cet enregistrement dans son cache. Le TTL est défini au niveau du SOA (et donc pour tous les RR situés dedans), mais peut, optionnellement, être défini pour chaque RR. Le délai de validité du cache est le maximum des valeurs <i>ttl</i> et <i>minimum</i> .			
IN	Désigne la classe d’adresse, dans notre cas l’Internet (on trouve aussi CH pour les adresses Chaos).			
RR	Nom de la ressource (SOA, NS, A, LNAME, NX, PTR, HINFO, etc.).			
nom	[ttl]	[IN]	SOA	name e-mail serial refresh retry expire minimum
(Start Of Authority). Indique le nom de la zone pour laquelle le serveur a autorité (secondaire ou primaire), ainsi que les paramètres de mise à jour de la base de données.				
name	Désigne le serveur primaire de la zone.			
e-mail	Désigne l’adresse e-mail de la personne responsable de la zone.			
serial	Désigne le numéro de version du SOA.			
refresh	Indique le nombre de secondes au bout duquel le serveur secondaire doit redemander le SOA au serveur primaire.			
retry	Indique le nombre de secondes qui s’écoule entre deux tentatives de téléchargement du SOA par le serveur secondaire.			
expire	Indique le nombre de secondes au bout duquel le SOA ne sera plus valable après le délai indiqué par <i>refresh</i> . Au-delà de ce délai, le serveur secondaire ne doit plus répondre aux requêtes concernant ce SOA.			
minimum	Indique la valeur minimale du TTL des RR de cette zone.			

...

LA BASE DE DONNÉES DU DNS (SUITE)

[nom] [ttl] (Name Server)	IN	NS	nom_du serveur_DNS	Indique le serveur (primaire ou secondaire) qui a autorité sur la zone.
[nom] [ttl] (Address)	IN	A	adresse_IP	Indique l'adresse IP qui correspond au nom (serveur, routeur, etc.) demandé dans une requête.
alias [ttl] (Canonical Name)	IN	CNAME	nom_principal_de_l'objet	Indique l'alias d'un objet.
nom [ttl] (Mail eXchanger)	IN	MX	priorité nom_du_MTA	Indique le nom du MTA acheminant les messages à destination du domaine. 65535 = priorité basse, 1 = priorité haute.
inverse [ttl] (Pointer) inverse	IN	PTR	nom_machine	Indique le nom qui correspond à l'adresse IP demandée dans une requête. Désigne l'adresse IP inverse, suivie de in-addr.arpa.
[nom] [ttl] (Host Information)	IN	HINFO	matériel système	Champ d'information concernant l'objet (type de machine et d'OS).
[nom] [ttl]	IN	AAAA	adresse_IPv6	Indique l'adresse IPv6 qui correspond au nom demandé dans une requête (RFC 1886).

Les noms de serveurs indiqués dans les RR ne doivent pas être des alias.

Il existe d'autres enregistrements peu ou pas utilisés ou encore à l'état expérimental : WKS, TXT (RFC 1035), AFSDB, RP, X25, ISDN, RT (RFC 1183), NSAP (RFC 1706), GPOS (RFC 1712), SRV (RFC 2052) et KX (RFC 2230).

Activer la résolution de nom

Généralement, l'administrateur commence par remplir la base de données d'enregistrements "A". Ce type d'enregistrement permet aux utilisateurs d'obtenir l'adresse IP correspondant au nom, qui leur est plus familier. C'est toute l'utilité du DNS.

Nouvel hôte

Entrez le nom du nouvel hôte pour info.

Zone "info."

Ajouter un hôte

Terminé

Nom d'hôte : nt003

Adresse IP d'hôte : 10 . 0 . 0 . 110

☐ Créer un enregistrement PTR associé

Création d'un enregistrement "A".

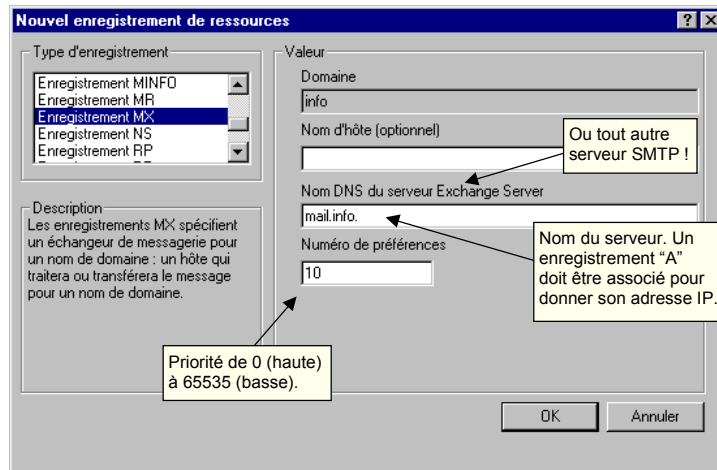
Cela a pour effet de créer l'enregistrement suivant :

```
nt003 IN      A      10.0.0.110
```

Désormais, le serveur DNS enverra au client l'adresse IP indiquée ci-dessus, à toute requête concernant ce nom.

Activer le routage de la messagerie

Un autre grand utilisateur du DNS est la messagerie SMTP (voir chapitre 3). Chaque MTA (*Message Transfer Agent*) s'appuie en effet sur les enregistrements "MX" pour router les messages vers le prochain MTA.



En définitive, deux enregistrements doivent être créés :

info.	IN	MX	10	mail.info.	← Ces deux notations sont équivalentes.
@	IN	MX	10	mail	
mail	IN	A	10.0.0.100		

Ainsi, tous les messages à destination de xxx@info seront routés vers le serveur "mail" situé dans le domaine "info".

Avec le mécanisme des priorités, il est possible de définir des serveurs de secours au cas où le MTA principal serait surchargé ou en panne. Il suffit pour cela de créer un autre enregistrement "MX" de priorité plus basse que celui créé précédemment :

```
info. IN      MX      20      mail2.info.
mail2 IN      A       10.0.0.101
```

Enfin, il est également possible de définir un serveur "poubelle" recueillant tous les messages à destination de domaines inconnus :

```
*      IN      MX      100     mail9
```

Du bon usage des alias

L'utilisation d'un alias permet de dissocier la fonction (web, messagerie) de la nature des serveurs (nt001, unix002) ou de leur localisation dans un domaine ou un autre :

```
www      IN      CNAME  nt003.info.
pluton  IN      CNAME  ux001.paris.compta.
```

Grâce à l'utilisation des alias, le changement des noms pluton et www suite à un déménagement sera plus facile à opérer que celui de nt003 qui est référencé par des enregistrements "A".

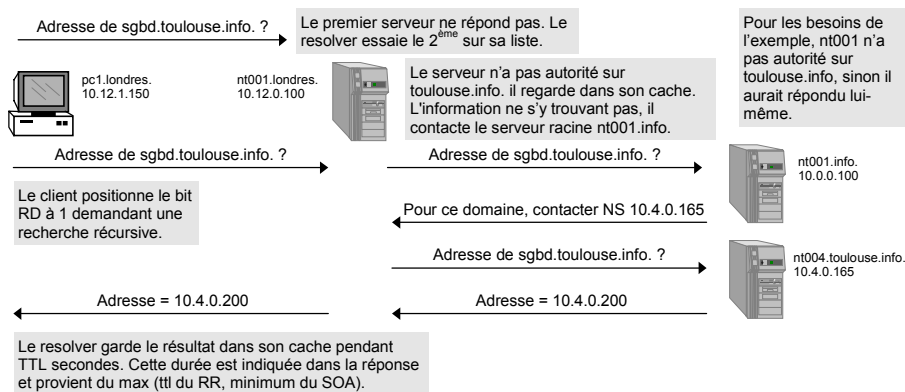
Par exemple, on doit déplacer l'application intranet sur une nouvelle machine plus puissante. Il suffit de modifier le CNAME comme suit :

```
www      IN      CNAME  nt004.toulouse.info.
```

Configurer un serveur racine

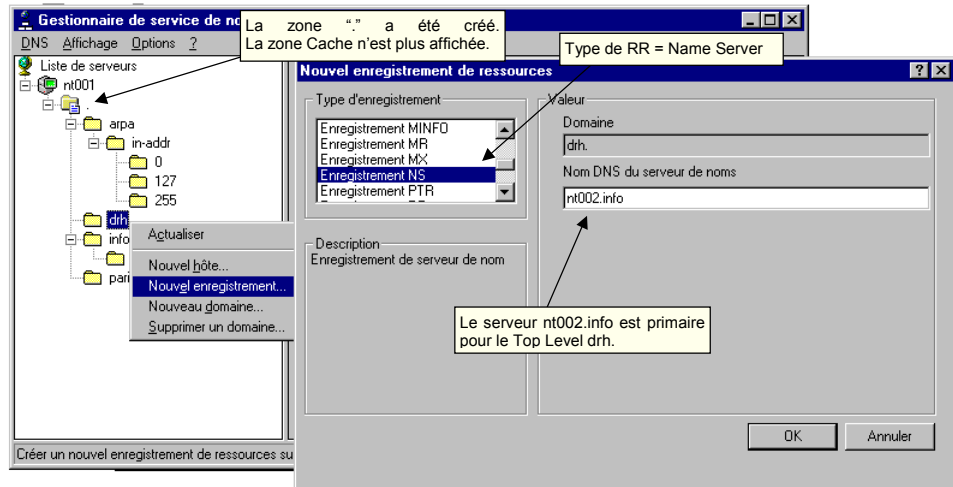
Les serveurs racines ont tout simplement autorité pour la zone '.', c'est-à-dire sur l'ensemble de l'arbre DNS. Leur rôle est d'indiquer aux serveurs ne pouvant répondre aux requêtes de leurs clients, l'adresse du serveur pouvant les aider.

Figure 17-3.
*Recherche DNS
récursive.*



La recherche peut également être effectuée de manière itérative. Dans ce cas, le serveur DNS indique au client le nom du serveur DNS qui peut l'aider, et il revient au client d'effectuer lui-même la recherche.

Dans notre cas, nous commençons par déclarer un serveur primaire pour la zone “.” (ne pas choisir le nom de fichier “.”, mais, par exemple, le nom “racine.dns”). Puis, nous déclarons les serveurs primaires des domaines Top Level.



Cela produit le fichier suivant :

```

.      IN      SOA      nt001.info.  jlm.nt001.info. (
                2      ; serial : 2ème version
                1800   ; refresh : toutes les 30 min.
                300    ; retry : toutes les 5 minutes
                604800 ; expire : 48 heures après
                86400) ; minimum : 24h dans le cache
      NS      nt001.info
      NS      nt004.toulouse.info

nt001.info      A      10.0.0.100
nt004.toulouse.info A    10.4.0.165

paris           NS     nt001.info
info            NS     nt001.info
drh             NS     nt002.info

nt002.info.     A      10.0.0.124

```

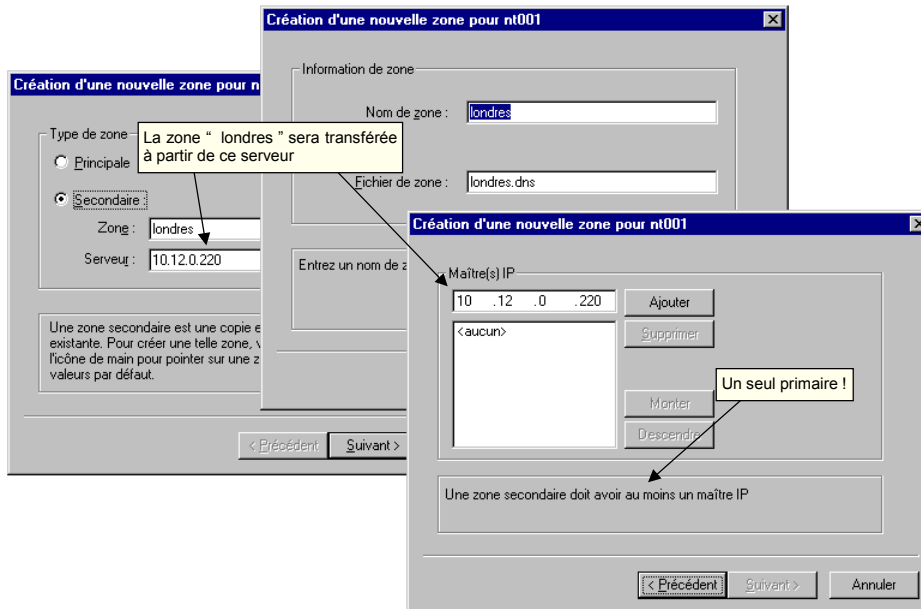
Liste des serveurs ayant
autorité sur la racine.

Liste des serveurs ayant
autorité sur les Top Level
paris, info et drh.

Configurer un serveur secondaire

Un serveur secondaire pour une zone est un serveur qui a autorité sur ladite zone, mais qui ne peut pas modifier la base de données correspondante. Pour cela, il demande auprès du serveur primaire **le transfert de la zone**, ce qui lui permet ensuite de répondre aux requêtes de la même manière qu'un serveur primaire.

Sur ce type de serveur, une zone est déclarée secondaire en indiquant l'adresse IP du serveur primaire à partir duquel transférer la zone.



Il est à noter que le volume de données représenté par un transfert de zone est peu important, d'autant plus si les serveurs supportent une mise à jour incrémentale (RFC 1995).

Configurer un serveur cache

Tous les serveurs DNS conservent en mémoire cache les réponses des requêtes précédentes. Le serveur cache joue le même rôle, mais présente la particularité de n'avoir aucune autorité sur une quelconque zone : le serveur n'est ni primaire ni secondaire.

La configuration d'un serveur cache se résume donc à renseigner uniquement le fichier "cache.dns" qui liste les serveurs racines.

La durée de validité des enregistrements dans le cache est déterminée par le maximum des deux valeurs suivantes : le paramètre "minimum TTL" du SOA et le TTL de chaque enregistrement si celui-ci est spécifié.

Déléguer l'autorité à un autre serveur

Pour des questions d'organisation, il est maintenant opportun de déléguer la gestion de la zone "toulouse.info" aux exploitants de ce site. Pour cela, la procédure est la suivante :

1. Configurer le serveur de Toulouse en secondaire de "toulouse.info", afin de transférer cette zone.
2. Configurer ensuite le serveur de Paris en secondaire pour la zone "toulouse.info".
3. Basculer enfin le serveur de Toulouse en primaire pour la zone "toulouse.info".
4. Modifier les enregistrements NS sur les différents serveurs concernés, de manière à pointer sur le nouveau serveur primaire.
5. Si l'ancien serveur primaire ne doit pas faire office de serveur secondaire du domaine "toulouse.info", supprimer toute référence à cette zone.

Sous le gestionnaire DNS de Windows, il faut, en plus, créer un domaine "toulouse.info", de manière à extraire les données du fichier "info.dns" dans un autre fichier, "toulouse.info.dns", par exemple.

Dans le serveur primaire de "info", la zone déléguée apparaîtra alors comme suit :

```
;
; Delegated sub-zone:  toulouse.info.
;
toulouse                NS      nt001
; End delegation
```

Les domaines de résolution inverse

Un domaine spécifique, appelé domaine de résolution inverse et noté **in-addr.arpa**, est utilisé pour trouver un nom à partir d'une adresse IP. Un domaine correspondant au réseau 10.4.0.0 sera ainsi noté 4.10.in-addr.arpa. (notation inverse de l'adresse IP).

Les domaines de résolution inverse se manipulent de manière identique aux domaines de noms. Chacun des 4 nombres de l'adresse IP est considéré comme étant un domaine qui peut être délégué. Un serveur peut ainsi être primaire pour le domaine 10.in-addr.arpa, et déléguer son autorité pour 15.10.in-addr.arpa.

La résolution inverse permet :

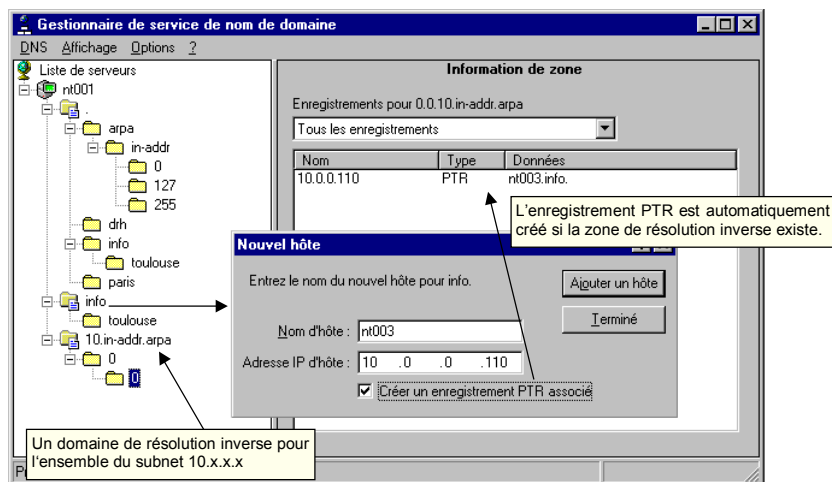
- De découvrir les routeurs d'un sous-réseau IP. Le client reçoit le nom des routeurs et peut ensuite lancer une requête pour en connaître les adresses IP. Pour des questions de sécurité, il est donc déconseillé d'utiliser cette facilité.
- À un serveur de s'assurer qu'un client qui se présente avec une adresse IP appartient bien à un domaine autorisé et/ou est bien celui qu'il prétend être. Certains serveurs FTP de l'Internet nécessitent que les clients soient référencés dans la base DNS sous forme d'enregistrement PTR.

- De référencer les noms de réseaux permettant à une station d'administration d'afficher les noms au lieu des adresses IP.
- Au programme **nslookup** (voir plus loin) de fonctionner correctement. En effet, le serveur à partir duquel la commande est lancée doit être référencé.

Une arborescence similaire doit donc être créée. Et, là encore, le DNS permet toutes sortes de fantaisies.

Par exemple, il n'y a pas obligatoirement de correspondance entre domaines de noms et domaines de résolution inverse. Les serveurs primaires pour un réseau peuvent ainsi ne pas l'être pour les domaines qui le contiennent. La création automatique de PTR est alors rendue difficile.

Il est donc conseillé de limiter le nombre de ce type de zones à une par subnet IP principal, le réseau 10 dans notre cas.

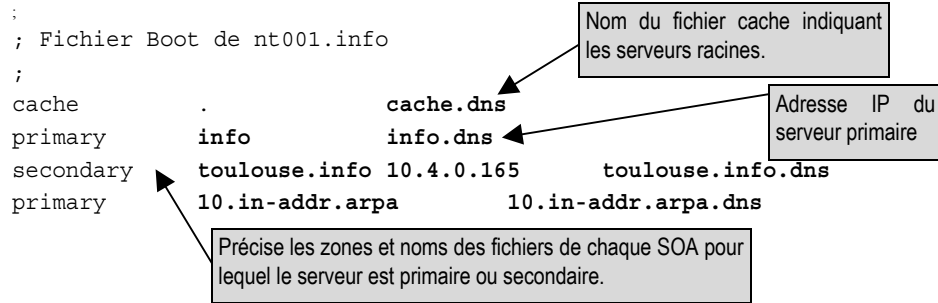


Le fichier d'initialisation

Le serveur DNS de Microsoft peut fonctionner sans l'interface graphique, à la manière des premiers serveurs DNS sous Unix. Le système fonctionne alors uniquement à partir des fichiers ASCII situés dans le répertoire `\Winnt\System32\Dns`.

Pour passer en mode texte, il faut supprimer la variable *EnableRegistryBoot* de la base des registres Windows au niveau de la clé "HKEY_LOCAL_MACHINE\System\CurrentControlSet\Services\Dns\Parameters".

Un nouveau fichier doit alors être renseigné : il s'agit du fichier d'initialisation du DNS tel qu'utilisé sous Unix, appelé **boot** :



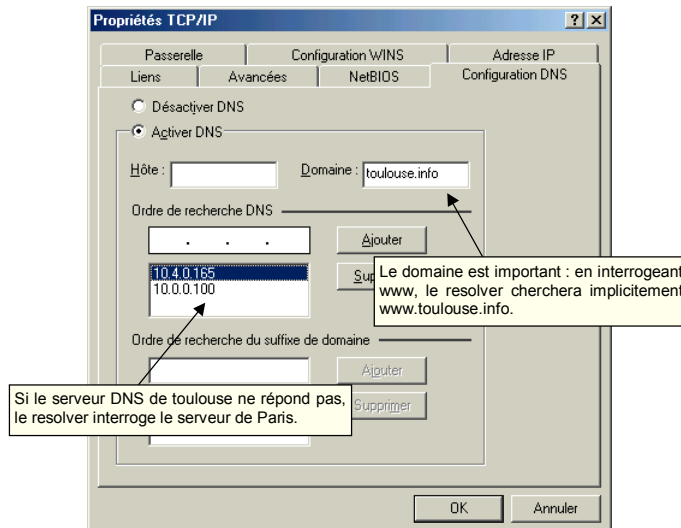
Configurer les clients DNS

Dans la terminologie DNS, un client est appelé **resolver**. Ce petit morceau de programme est directement sollicité par toutes les applications fonctionnant en réseau, telles que le navigateur web. Il se contente d'interroger des serveurs DNS.

Sa configuration consiste simplement à indiquer le domaine dans lequel il se trouve, ainsi que la liste des adresses IP des serveurs DNS susceptibles de répondre à ses requêtes.

Une fois de plus, le DNS offre beaucoup de souplesse :

- le client interroge le premier serveur de la liste puis, s'il ne répond pas, le second et ainsi de suite ;
- les serveurs peuvent être de type cache, primaire ou secondaire ;
- le client peut être situé dans un domaine différent des serveurs qu'il interroge.



L'utilisateur peut formuler une demande sur un **nom complet**, tel que “`www.info.`” (un nom terminé par un point). Aucun suffixe ne sera alors ajouté.

S'il lance une recherche sur un **nom relatif**, par exemple `www` (un nom qui n'est pas terminé par un point), le resolver le recherchera par défaut dans le même domaine que celui configuré. S'il reçoit une réponse négative, il ajoutera le suffixe `".info"`, puis un autre suffixe en fonction de la configuration.

Vérifier le fonctionnement du DNS

L'utilitaire de base pour tout administrateur DNS est le **nslookup** qui s'exécute à partir d'une fenêtre DOS de Windows NT :

```

>set d2
>set recurse
>www.3com.com.
> Serveur:  nt001.info.
Address:  10.0.0.100

-----
SendRequest(), len 30

    HEADER:
opcode = QUERY, id = 5, rcode = NOERROR
header flags:  query, want recursion
questions = 1,  answers = 0,  authority records = 0,  additional = 0

    QUESTIONS:
      www.3com.com, type = A, class = IN

-----
-----
La question porte sur un
enregistrement "A".

Got answer (301 bytes):

    HEADER:
opcode = QUERY, id = 5, rcode = NOERROR
header flags:  response, want recursion, recursion avail.
questions = 1,  answers = 1,  authority records = 6,  additional = 6
  
```

Se positionne en mode debug afin d'afficher le détail des échanges.

Une recherche récursive est demandée

Nom et adresse du serveur DNS de rattachement.

On demande l'adresse IP du serveur `www` situé dans le domaine `3com.com`.

Dans sa réponse, le serveur DNS indique qu'il supporte la recherche récursive.

Nombre de RR par section

```

QUESTIONS: ←
    www.3com.com, type = A, class = IN
ANSWERS: ←
->  www.3com.com
    type = A, class = IN, dlen = 4
    ttl = 3645 (1 hour 45 secs)
AUTHORITY RECORDS: ←
->  3COM.COM
    type = NS, class = IN, dlen = 9
    nameserver = FOUR11.3COM.COM
    ttl = 38819 (10 hours 46 mins 59 secs)
..... suivent 5 autres RR .....
ADDITIONAL RECORDS: ←
->  FOUR11.3COM.COM
    type = A, class = IN, dlen = 4
    internet address = 129.213.128.98
    ttl = 92825 (1 day 1 hour 47 mins 5 secs)
..... suivent 5 autres RR .....

```

La réponse est structurée en 4 sections :
question, réponse, enregistrements NS et
les autres enregistrements.

Par défaut, l'enregistrement " A " est demandé, mais il est possible de demander d'autres types d'enregistrements, tels que le " SOA " :

```

>set d2
>set recurse
>set querytype=soa
>www.3com.com.
> Serveur:  ns1.club-internet.fr
Address:  194.117.200.10

```

SendRequest(), len 26

HEADER:

```

opcode = QUERY, id = 5, rcode = NOERROR
header flags:  query, want recursion
questions = 1,  answers = 0,  authority records = 0,  additional = 0

```

QUESTIONS:

```

    3com.com, type = SOA, class = IN

```


Got answer (320 bytes):

HEADER:

opcode = QUERY, id = 5, rcode = NOERROR
header flags: response, want recursion, recursion avail.
questions = 1, answers = 1, authority records = 6, additional = 6

QUESTIONS:

3com.com, type = SOA, class = IN

ANSWERS:

```
-> 3com.com
    type = SOA, class = IN, dlen = 42
    ttl = 68436 (19 hours 36 secs)
    primary name server = four11.3com.com
    responsible mail addr = hostmaster.3com.com
    serial = 1998090100
    refresh = 10800 (3 hours)
    retry = 3600 (1 hour)
    expire = 604800 (7 days)
    default TTL = 10800 (3 hours)
```

On retrouve les paramètres du SOA
tels que l'administrateur de 3com les
a définis.

AUTHORITY RECORDS:

```
-> 3com.com
    type = NS, class = IN, dlen = 2
    nameserver = four11.3com.com
    ttl = 38922 (10 hours 48 mins 42 secs)
..... etc. ....
```

Annexes

Normes et standards

Le câblage

Normes CEN relatives au câblage

Référence	Objet
EN 55022	Compatibilité électromagnétique pour les équipements informatiques
EN 55024	Protection contre les champs électromagnétiques (dérivé de IEC 801)
EN 50081-1 EN 50082-1	Compatibilité électromagnétique des composants électroniques et des câbles
EN 50167 EN 50168 EN 50169	Caractéristiques et performances des câbles FTP
EN 50173	Spécifications des performances des systèmes de câblage (dérivé de l'ISO/IEC DIS 11801)
EN 186000-1	Spécifications génériques pour les câbles et connecteurs fibres optiques
EN187000	Spécifications générales des câbles fibres optiques
EN 188000	Spécifications générales des fibres optiques

La norme AWG (*American Wire Gauge*) définit, quant à elle, la section des conducteurs électriques dans un câble.

Normes EIA/TIA relatives au câblage

Référence	Objet
EIA/TIA-569	Spécifications générales des infrastructures d'immeuble (espace alloué, chemins de câbles...)
EIA/TIA-568	Spécifications générales des systèmes de câblage
EIA/TIA-586	Définition des 5 catégories de câbles UTP (<i>Unshielded Twisted Pair</i>)
EIA/TIA TSB 36	Spécifications des câbles de catégorie 5
EIA/TIA TSB 40	Spécifications des connecteurs de catégorie 5
EIA-455-48	Spécifications du cœur des fibres optiques 62.5/125
EIA-455-27 EIA-455-48	Spécifications des caractéristiques mécaniques du revêtement des fibres optiques 62.5/125
EIA-455-57	Spécifications de l'ouverture numérique des fibres optiques 62.5/125

Normes ITU-T relatives au câblage

Référence	Objet
ITU-T G651	Caractéristiques des câbles fibres optiques à gradient d'indice 50/125µm
ITU-T G652	Caractéristiques des câbles monomode

Les interfaces physiques

Avis de l'ITU-T relatifs aux interfaces physiques

Avis	Contenu de la spécification
V.6	Standardisation des vitesses de transmission synchrone sur les lignes spécialisées
V.11	Caractéristiques électriques des circuits d'échange à double courant symétrique jusqu'à 10 Mbit/s
V.24	Définition des échanges entre ETTD et ETCD et du format de la prise
V.28	Caractéristiques électriques pour les circuits d'échange double courant asymétrique
V.33	Modem 14,4 Kbit/s <i>full duplex</i> sur 4 fils
V.34	Modem 33,6 Kbit/s <i>full duplex</i> (V.FAST)
V.42	Correction d'erreurs LAPM (<i>Link Access Procedure for Modem</i>)
V.42 bis	Compressions de données
V.54	Équipement de test de modem en boucle
V.90	Modem à 56,4 Kbit/s

Normes EIA/TIA relatives aux interfaces physiques

Référence	Objet
EIA/TIA-232	Spécifications de l'interface RS-232 proches de celles de la norme V.24. Débit maximal de 64 Kbit/s.
EIA/TIA-449	Spécifications de l'interface RS-449. Amélioration de la norme RS-232. Débit maximal de 2 Mbit/s.
EIA-530	Spécifications de deux implémentations d'échange de la norme EIA-TIA-449 : mode balancé (RS-422) et mode non balancé (RS-423).

Avis de l'ITU-T relatifs aux échanges ETTD-ETCD

Avis	Contenu de la spécification
X.20	Interface entre un ETTD et un ETCD pour les services de transmission asynchrone
X.21	Interface entre un ETTD et un ETCD pour fonctionnement synchrone
X.24	Définitions des circuits d'échanges entre un ETTD et un ETCD

Les réseaux locaux

Normes IEEE relatives aux réseaux locaux

Référence	Désignation	Objet
802.1	<i>High Level Interface</i>	Traite des architectures (802.1 a), des ponts et du <i>spanning tree</i> (802.1d) et du <i>System Load Protocol</i> (802.1 e).
802.1p	<i>Traffic Class Expediting and Dynamic Multicast Filtering</i>	Gestion du flux et des priorités sur Ethernet
802.1q	VLAN <i>Virtual Bridged Local Area Networks</i>	Ajoute un entête de 4 octets (une étiquette ou encore <i>tag</i>) aux trames Ethernet définissant le numéro de réseau virtuel
802.2	LLC <i>Logical Link Control</i>	Spécifications de la sous-couche LLC du niveau 2 du modèle OSI (802.2c, f et h)
802.3	Ethernet CSMA/CD	Spécifications des réseaux Ethernet
802.3u	Ethernet 100bT	Spécifications du Fast Ethernet. Couche MII (<i>Media Independant Interface</i>), 100bTX, 100bT4...
802.3x	Full Duplex et contrôle de flux	Signal intercommutateurs émis pour arrêter le trafic lorsque la mémoire est saturée
802.3ab	Ethernet 1000bT	Spécifications du Gigabit Ethernet sur cuivre en paires torsadées
802.3z	Ethernet 1000bX	Spécifications du Gigabit Ethernet sur fibre optique
802.4	Réseaux Token-Bus	Spécifications des réseaux Token-Bus
802.5	Réseaux Token-Ring	Spécifications des réseaux Token-Ring
802.6	Réseaux MAN DQDB	Spécifications des réseaux métropolitains
802.7	Réseaux large bande	Groupe de travail BBTAG (<i>Broadband Technical Advisory Group</i>). Norme <i>Slotted Ring</i> .
802.8	Réseaux fibre optique	Groupe de travail FOTAG (<i>Fibre Optics Technical Advisory Group</i>)
802.9	Réseaux voix/données	IS LAN (<i>Integrated Services LAN</i>) Ethernet Isochrone - IsoEnet
802.10	Sécurité des réseaux	Méthodes d'accès entre les couches MAC et LLC (niveau 2) ainsi que pour la couche application (niveau 7) pour les données confidentielles
802.11	Réseaux sans fil	WLAN (<i>Wireless LAN</i>)
802.12	100bVG-AnyLAN	Spécifications des réseaux locaux à 100 Mbit/s avec DPMA (<i>Demand-Priority Access Method</i>)
802.14	CATV (<i>Cable-TV</i>)	Réseaux sur les câbles télévision CATV

La famille des protocoles TCP/IP

RFC relatives aux protocoles TCP/IP

Référence	Objet
791	Spécifications de IP (<i>Internet Protocol</i>) - MIL-STD-1777
792	Spécifications de ICMP (<i>Internet Control Message Protocol</i>)
793, 761, 675	Spécifications de TCP (<i>Transmission Control Protocol</i>) - MIL-STD-1778
768	Spécifications de UDP (<i>User Datagram Protocol</i>)
813	Algorithmes d'acquittement de TCP
815	Algorithmes de réassemblage des paquets TCP/IP
816	Mécanisme de dead gateway
919, 922	Diffusion des datagrammes Internet (<i>broadcast</i>)
917, 932, 936, 950	Spécifications et description du subnetting des adresses IP
1219	Subnetting variable
826	ARP (<i>Address Resolution Protocol</i>) sur Ethernet
903	RARP (<i>Reverse Address Resolution Protocol</i>)
1293	Inverse ARP
1027	Proxy ARP (pour les stations ne supportant pas les <i>subnets</i>)
1011	Description officielle des protocoles Internet
1108	Spécifications de IPSO (<i>IP Security Option</i>)
894	Encapsulation de IP dans une trame Ethernet V2
1042	Encapsulation de IP dans une trame Ethernet 802.3
1078	Multiplexage des ports des services TCP
1144	Compression des en-têtes TCP
1505	En-tête des messages Internet
1918	Subnets IP réservés à l'adressage privé
1001 - 1002	Netbios sur TCP/IP : concepts et spécifications
1356	Encapsulation dans X25
1434	DLSw - Encapsulation des trames SNA dans des paquets TCP/IP
1700	Liste des valeurs attribuées aux différents champs des protocoles de la famille TCP/IP par le IANA (<i>Internet Assigned Numbers Authority</i>)

Standards originaux du DOD (Department Of Defense) relatifs à TCP/IP

Référence	Objet
MIL-STD-1777	Spécifications de IP (<i>Internet Protocol</i>)
MIL-STD-1778	Spécifications de TCP (<i>Transport Control Protocol</i>)
MIL-STD-1780	Spécifications de FTP (<i>File Transfer Protocol</i>)
MIL-STD-1781	Spécifications de SMTP (<i>Simple Mail Transfer Protocol</i>)
MIL-STD-1782	Spécifications de Telnet

RFC relatives aux protocoles de routage IP

Référence	Objet
1256	Messages ICMP de découverte de routes
2328	OSPF (<i>Open Shortest Path First</i>) version 2
1245, 1246	Analyse du fonctionnement de OSPF
1771, 1772	BGP (<i>Border Gateway Protocol</i>)
827, 904, 911	EGP (<i>Exterior Gateway Protocol</i>)
1058, 1723	RIP (<i>Routing Information Protocol</i>)
950	Procédures de <i>subnetting</i>

RFC relatives aux applications utilisant TCP/IP

Référence	Objet
821, 822, 974	Spécifications de SMTP (<i>Simple Mail Transfer Protocol</i>)
1869	Extended SMTP
1777, 1778	Spécifications de LDAP v3 (<i>Lightweight Directory Access Protocol</i>)
1344, 1437	Spécifications de Mime (<i>Multipurpose Internet Mail Extensions</i>)
2045, 2046	Structure des messages Mime et des formats supportés
1939	Spécifications de POP 3 (<i>Post Office Protocol</i>)
2060	Spécifications de IMAP 4 (<i>Internet Message Access Protocol</i>)
854	Spécifications de Telnet
1205	Spécifications de l'émulation 5250 sur Telnet
1184, 1091, 1372	Spécifications de diverses options Telnet
1282	Spécifications de Rlogin (<i>remote login</i>)
959	Spécifications de FTP (<i>File Transfert Protocol</i>)
1350	Spécifications de TFTP (<i>Trivial File Transfer Protocol</i>)

Référence	Objet
1094, 1813	Spécifications de NFS (<i>Network File System</i>)
1945	Spécifications de HTTP 1.0 (<i>Hyper Text Transfert Protocol</i>)
2068, 2069, 2109, 2145	Spécifications de HTTP 1.1 (<i>Hyper Text Transfert Protocol</i>)
1630, 1738	Spécifications des URL (<i>Uniform Resource Locators</i>)
1034, 1035	Concept, spécifications et implémentation du DNS (<i>Domain Name System</i>)
1982, 1995, 1996, 2136, 2137, 2181, 2308	Mises à jour de DNS
1183, 1706, 1712, 2052, 2230	Extensions DNS expérimentales
2535	Extensions DNS sécurisé
1591, 1912	Implémentations et bonnes pratiques du DNS
906, 951, 1542	Spécifications de BootP (<i>Bootstrap Protocol</i>)
1534	Interopérabilité entre DHCP et BootP
2131	Spécifications de DHCP (<i>Dynamic Host Configuration Protocol</i>)
2132, 2224	Options DHCP

RFC relatives à IP sur Frame-Relay

Référence	Objet
2427	Transport de IP dans Frame-Relay (NLPID (<i>Network Level Protocol ID</i>) / SNAP (<i>Sub Network Access Protocol</i>))

RFC relatives à IP sur ATM

Référence	Objet
1332	Base de travail pour IP sur ATM (<i>Asynchronous Transfert Mode</i>)
1483	Utilisation de la couche AAL-5 pour l'encapsulation des protocoles IP dans un réseau ATM. Spécifications DXI (<i>Data eXchange Interface</i>)
2225	Spécifications du routage IP et de la résolution d'adresse ARP sur les réseaux ATM (Classical IP)
1626	Taille des paquets IP pour ATM AAL-5
1629	Guide pour l'allocation des adresses NSAP (<i>Network Service Access Point</i>) au sein de l'Internet
1680	Support des Services ATM pour IPv6
1755	Signalisation pour IP sur ATM
2022	Multicast sur UNI 3.0/3.1
2149	Spécifications de MARS (<i>Multicast Server Architectures for MARS-based ATM multicasting</i>)

RFC relatives à PPP

Référence	Objet
1661, 1662, 1663	Spécifications du protocole PPP
1332	Spécifications de la couche IPCP (<i>IP Control Protocol</i>)
1552	Spécifications de la couche IPXCP (<i>IPX Control Protocol</i>)
1570	Spécifications de la couche LCP (<i>Link Control Protocol</i>)
1989	Spécifications de la couche LQM (<i>Link Quality Monitoring</i>)
1990	Spécifications de PPP MP (<i>PPP Multilink Protocol</i>)
1321	Spécifications de l'algorithme MD5 (<i>Message Digest 5</i>)
1993	Algorithme de compression FZA pour PPP
1994	Spécifications des protocoles CHAP (<i>Challenge-Handshake Authentication Protocol</i>)

RFC relatives à SNMP

Référence	Objet
1089, 1157	Spécifications de SNMP (<i>Simple Network Management Protocol</i>)
1303	Description conventionnelle des agents SNMP
1352	Protocoles de sécurité SNMP
2571	Architecture de l'administration avec SNMP
2572	Traitement des messages SNMP
1418, 1419, 1420	SNMP sur OSI, AppleTalk et IPX
1115	Structure de la MIB pour les protocoles TCP/IP
1212	Spécifications génériques de la MIB (<i>Management Information Base</i>)
1213	Spécifications de la MIB-II pour TCP/IP
2863	Le groupe « Interface » de la MIB
2011, 2012, 2013	MIB-II pour SNMP V2 (mises à jour du RFC 1213)
1231	Spécifications de la MIB Token Ring
1381	Spécifications de la MIB X25
1382	Spécifications de la MIB HDLC LAP-B
1398	Spécifications de la MIB Ethernet
1512, 1285	Spécifications de la MIB FDDI
1513, 1271	Spécifications de la MIB RMON (<i>Remote Monitoring</i>)
1559	Définition de la MIB Decnet Phase IV

Normes ISO et équivalents ITU-T relatifs à la syntaxe ASN.1

Référence ISO	Référence ITU-T	Objet
8824	X.208	Spécifications du langage ASN.1 (<i>Abstract Syntax Notation 1</i>)
8825	X.209	Spécifications des règles de codage BER (<i>Basic Encoding Rules</i>)

RFC relatives à IPv6

Référence	Objet
1883 et 2147	Spécifications de IPv6 (IPng - <i>IP Next Generation</i>)
2133	Extension des sockets pour IPv6
2073, 1887	Format des adresses <i>unicast</i> ,
1924, 1971, 1897, 1884, 1881	Adressage IPv6 (configuration, allocations)
1809	Utilisation du champ de contrôle de flux
2019	IPv6 sur FDDI
1972	IPv6 sur Ethernet
1970	Protocole de découverte des nœuds IPv6
1933	Procédures de transmission des paquets IPv6
1888	NSAP OSI pour IPv6
1885	ICMPv6 (<i>Internet Control Message Protocol</i>)
1886	Extension du DNS pour le support de Ipv6
2080, 2081	RIPng (<i>Routing Information Protocol next generation</i>) pour IPv6
de 1667 à 1688, 1550, 1705, etc.	Discussions sur IPng (<i>Internet Protocol next generation</i>)

Le multimédia sur IP (VoIP)

RFC relatives à la voix sur IP

Référence	Objet
2543	SIP (<i>Session Initiation Protocol</i>)
2976	Extension de SIP : méthode INFO
1889	RTP (<i>Real-time Transport Protocol</i>) et RTCP (<i>Real-time Transport Control Protocol</i>)
1890	Définition des profils pour les conférences audio et vidéo
2032	Transport des codec vidéo H.261
2190	Transport des codec vidéo H.263
2198	Transport des codec audio

Avis de l'ITU-T relatifs à la voix sur IP

Référence	Objet
H.323	Systèmes de communications Multimedia basé sur des paquets
H.323 – Annexe A	Messages H.245 utilisés par H.323
H.323 – Annexe B	Procédures pour les codec vidéo
H.323 – Annexe C	H.323 sur ATM AAL5
H.323 – Annexe D	Fax temps réel sur H.323
H.323 – Annexe E	Transport de la signalisation d'appel
H.323 – Annexe G	Spécifications des terminaux H.323
H.323 – Annexe H	Mobilité
H.323 – Annexe I	Opération sur des réseaux avec une basse qualité de service
H.323 – Annexe J	Terminaux sécurisés
H.323 – Annexe K	Interface avec http
H.323 – Annexe L	Gestion des communications un serveur (<i>feature server</i>)
H.323 – Annexe M	Tunneling Q.SIG
H.323 – Annexe N	Qualité de service
H.225	Protocole de signalisation d'appel (RAS et Q.931)
H.235	Sécurité
H.245	Protocole de contrôle pour les communications multimédias
Q.931	UNI niveau 3 pour le contrôle d'appel
H.450	Protocole générique pour les services complémentaires

Référence	Objet
H.261	Codec vidéo
H.263	Codec vidéo
G.711, G.722, G.723	Codec audio
G.728, G.729	Codec audio
T.120	Protocole de données pour les conférences multimédias

RFC relatives à la qualité de service

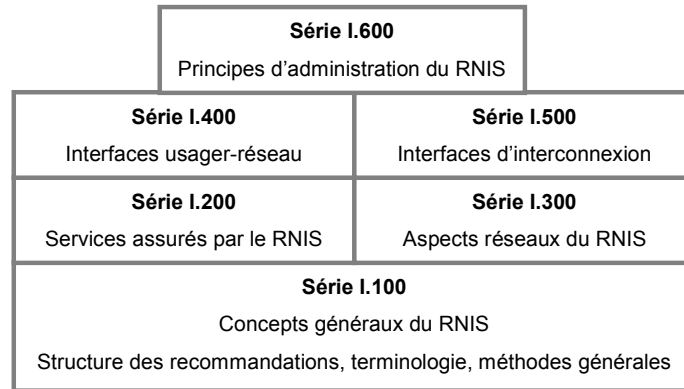
Référence	Objet
791, 1349	Champs TOS (<i>Type Of Service</i>) et Precedence du paquet IP
2474, 2475	DiffServ (<i>Differentiated Service</i>)
2873	Compatibilité du champ TOS avec le Code Point de DiffServ
1633	IntServ (<i>Integrated Service</i>)
2205	RSVP (<i>Resource reSerVation Protocol</i>)
2206-2209	Extensions MIB et IP-SEC et mise en œuvre
2210	Utilisation de RSVP dans IntServ
2211	IntServ – Classe de service par contrôle de charge
2212	IntServ – Classe de service par garantie de service
2213	IntServ – MIB de l'intégration de service
2214	IntServ – MIB de l'intégration de service - Extension à la garantie de service
2215	IntServ – Définition des paramètres permettant de calculer une QoS
2216	IntServ – Définitions de la QoS pour les éléments du réseau (définit un cadre général sur lequel s'appuient les RFC 2211 et 2212)

RFC relatives au routage multicast

Référence	Objet
1112	Support multicast par les piles TCP/IP
2236	IGMP (<i>Internet Group Membership Protocol</i>)
1075	DVMRP (<i>Distance Vector Multicast Routing Protocol</i>)
1584	MOSPF (<i>Multicast Open Shortest Path First</i>)
2362	PIM-SM (<i>Protocol Independent Multicast</i>)

Les réseaux RNIS

Organisation et nomenclature des normes relatives aux réseaux numériques à intégration de services



Le terme RNIS (réseau numérique à intégration de services) ou ISDN (*Integrated Services Digital Network*) désigne le réseau connu des utilisateurs (c'est-à-dire l'accès de base T0 et l'accès primaire T2) et, plus généralement, l'ensemble des réseaux numériques comme l'ATM, le relais de trames, SDH, etc. Tous ces protocoles sont fédérés par un ensemble de normes communes. L'ATM est ainsi référencé sous le nom de RNIS large bande (*Broad-band ISDN*).

	RNIS		ATM		Relais de Trames	
Couche 3	I.450 - Q.930 I.451 - Q.931 I.452 - Q.932	UNI Signalisation Procédures	Q.933 - Q.2931	UNI	Q.933 - Q.2931	UNI
Couche 2	I.440 - Q.920 I.441 - Q.921	UNI Lap-D	I.362, I.363 I.361 I.371	AAL ATM Congestion	Q.921, Q.922 I.370	Couche liaison Congestion
Couche 1	I.420, I.430 I.421, I.431	Accès de base Accès primaire	I.432	PHY	I.430 I.431	Accès de base Accès primaire

Série I.100 : Concepts généraux du RNIS

Avis	Contenu des spécifications
I.110	Structure générale des recommandations de la série I
I.111	Relations avec les autres recommandations
I.112	Vocabulaire RNIS (Réseau Numérique à Intégration de service)
I.120	Description du RNIS (ISDN - Integrated Services Digital Network)
I.130 et I.140	Services du RNIS et fonctionnalités du réseau (méthode et technique)

Série I.200 : Services assurés par le RNIS

Avis	Contenu des spécifications
I.200	Description de la série I.200
I.210	Généralités sur les services offerts par le RNIS
I.211	Définition des services réseaux (au niveau du SSTM)
I.212	Description des téléseuices
I.220	Description dynamique commune des services de télécommunications
I.221	Caractéristiques des services spécifiques courants
I.230	Définition des catégories des services support
I.231	Description des catégories des services support en mode circuit
I.232	Description des catégories des services support en mode paquet
I.240 et I.241	Définition et description des téléseuices offerts par le RNIS
I.250	Définition des compléments de service
I.251 à I.257	Description des compléments de service (identification de l'appelant, conférence, facturation, transfert d'appel...)

Série I.300 : Aspects réseaux du RNIS

Avis	Contenu de la spécification
I.310	Principes fonctionnels du réseau
I.320	Modèle de référence des protocoles
I.324	Architecture du RNIS
I.325	Configurations de référence pour les types de connexion du RNIS
I.326	Configurations de référence pour les besoins de ressources
I.330	Principes de la numérotation et de l'adressage

Avis	Contenu de la spécification
I.331	Plan de numérotation E.164
I.332	Principes d'adressage pour l'interconnexion de plusieurs RNIS
I.333	Sélection du terminal dans le RNIS
I.334	Principes d'adressage et de sous-adressage par rapport au modèle OSI
I.335	Principes de routage
I.340	Interconnexion des RNIS (attributs et topologies)
I.350	Objectifs de performance pour le réseau téléphonique et le RNIS
I.351	Objectifs de performance au point de référence T
I.352	Objectifs de performance en termes de délais de connexion

Série I.400 - Interfaces usager-réseau

Avis	Contenu de la spécification
I.410	Généralités et principes de la UNI (User Network Interface) : indique les types d'interfaces, le paramétrage (débit, codage...), la vérification de compatibilités entre terminaux, etc.
I.411	Configurations de référence : définitions de la TNR, TNA, etc.
Couche physique du RNIS (couche 1)	
I.412	Structure d'interface et possibilités d'accès Définitions des types de canaux : 2B+D Canal H0 à 384 Kbits/s Canal H11 à 1 536 Kbits/s (1 544 Kbits/s en accès primaire) Canal H12 à 1 920 Kbits/s (2 040 Kbits/s en accès primaire)
I.420	Structure des canaux et interface d'accès de base
I.421	Structure des canaux et interface d'accès primaire
I.430	Spécifications de la couche physique de l'accès de base. Interface usager-réseau. Points de référence S et T
I.431	Spécifications de la couche physique de l'accès primaire
Couche liaison du RNIS (couche 2)	
I.440 - Q.920	Généralités sur la signalisation UNI (User Network Interface)
I.441 - Q.921	Spécifications de la couche 2 : LAP-D, signalisation DASS (Digital Access Signaling System)
Couche réseau du RNIS (niveau 3)	
I.450 - Q.930	Généralités sur la signalisation UNI
I.451 - Q.931	Spécifications de la couche réseau signalisation des services
I.452 - Q.932	Procédures pour le contrôle des services complémentaires

Avis	Contenu de la spécification
Support d'autres interfaces par le RNIS	
I.460	Adaptation de débit au multiplexage et support des interfaces existantes
I.461 - X.30	Support des interfaces X21, X21 bis et X20 bis
I.462 - X.31	Support des terminaux en mode paquet
I.463 - V.110	Support des ETTD série V
I.464	Multiplexage, adaptation de débit et support des interfaces existantes pour des possibilités réduites de transfert à 64 Kbit/s
I.465 - V.120	Support des ETTD dotés d'interfaces de la série V et multiplexage statistique
I.470	Relations avec les fonctions du terminal RNIS

Série I.500 : Interfaces d'interconnexion du RNIS

Avis	Contenu de la spécification
I.500	Structure générale de l'interconnexion de réseaux RNIS
I.510	Définition et principes généraux de l'interconnexion
I.511	Interface physique de l'interconnexion
I.515	Échanges de paramètres
I.520	Spécifications générales de l'interconnexion des réseaux RNIS
I.530	Interconnexion avec un réseau public
I.540 - X.321	Spécifications générales pour l'interconnexion avec un réseau public à commutation de circuits
I.550 - X.325	Spécifications générales pour l'interconnexion avec un réseau public à commutation de paquets
I.560 - U.202	Support du service télex par le RNIS

Série I.600 : Administration du RNIS

Avis	Contenu de la spécification
I.601	Principes généraux de l'administration
I.602	Application des principes à l'installation d'un abonné
I.603	Application des principes à l'accès de base
I.604	Application des principes à l'accès primaire
I.605	Application des principes aux accès multiplexés

Avis de l'ITU-T relatifs aux réseaux ATM

Référence	Objet
I.113	Vocabulaire relatif au B-ISDN (Broadband-Integrated Services Digital Network)
I.121	Généralités sur les aspects large bande du B-ISDN
I.150 et I.351	Caractéristiques fonctionnelles ATM du B-ISDN
I.211	Présentation générale des services du B-ISDN
I.311 - I.312	Présentation générale du réseau B-ISDN
I.321	Modèle de référence du protocole B-ISDN et ses applications
I.327	Architecture fonctionnelle du B-ISDN
I.353	Définition des événements et paramètres permettant de mesurer les performances de la commutation de cellules
I.356	Performances de la commutation de cellules
I.361	Spécifications de la couche ATM : format des cellules
I.362	Description fonctionnelle de la couche AAL (ATM Adaptation Layer)
I.363	Spécifications de la couche AAL
I.371	Contrôle de flux et gestion des congestions
I.413	Spécifications de l'UNI (User Network Interface)
I.432	Spécifications de la couche physique (niveau 1) de l'UNI, de l'en-tête HEC et du mécanisme cell delineation
I.600	Application des principes de gestion des abonnés B-ISDN
I.610	Principes de gestion du B-ISDN (couche OAM - Operations, Administration and Maintenance)
Q.93B - G.771	Spécifications du contrôle des services et des appels (UNI dérivée de la norme Q.931).
Q.2931 - Q.933	Couche 3 UNI (version 3.0) de l'ATM Forum

Avis de l'ITU-T et équivalents ANSI relatifs au relais de trames

Référence ITU-T	Référence ANSI	Objet
I.233	T1.606	Description des services du relais de trames
I.370	T1.606	Gestion des congestions
I.372		Besoins pour les interfaces NNI
I.555		Relais de Trames sur ATM
Q.921		Description de la couche liaison pour la signalisation
Q.922	T1.618	Principes du protocole, spécifications de la couche liaison
Q.933	T1.617	Spécifications de la signalisation UNI de niveau 3

Avis de l'ITU-T relatifs aux systèmes de transmission numérique MIC

Référence	Objet
G.703	Caractéristiques physiques et électriques des interfaces
G.704	Caractéristiques fonctionnelles des jonctions
G.711	Modulation par impulsions codées (MIC) des fréquences vocales
G.721	ADPCM (Activity Detection Pulse Code Modulation) : codage MIC différentiel adaptatif (MIC-DA) à 32 Kbit/s
G.722	Codage des fréquences audio à 7 KHz sur 64 Kbit/s
G.725	Aspects relatifs à l'utilisation du codage audio à 7KHz sur 64 Kbit/s
G.726	ADPCM avec compression à 16k, 24k ou 32k
G.728	LDCELP : compression de la voix à 16k
G.729	CS-ACELP : compression de la voix de 4,8k à 16k
G.732	Caractéristiques des équipements de multiplexage MIC primaires à 2 048 Kbit/s.

Avis de l'ITU-T relatifs aux réseaux SDH

Référence	Objet
G.702	Spécifications des débits de la hiérarchie numérique
G.703	Caractéristiques physiques et électriques des interfaces MIC
G.704	Structure des trames des niveaux primaires et secondaires
G.706	Spécifications du CRC (Cyclic Redundancy Code) et de l'alignement des trames
G.707	Spécifications des débits de la hiérarchie numérique synchrone
G.708	Spécifications des interfaces des nœuds du réseau
G.709	Structure de multiplexage synchrone
G.733	Caractéristiques physiques et électriques des interfaces PCM (D2)
G.771 - Q.93B	Spécifications du contrôle des services et des appels (UNI)
G.774	Modèle d'administration pour les éléments du réseau
G.781	Aspects généraux des systèmes de transmission (multiplexeurs)
G.782	Caractéristiques générales des équipements de multiplexage
G.783	Caractéristiques générales des blocs fonctionnels des multiplexeurs
G.784	Aspects généraux des équipements d'administration
G.804	Format d'encodage des cellules ATM sur le support physique
G.821	Performance pour un réseau international
G.826	Paramètres et objectifs de performance pour les classes de trafic constant

A

ACR (*Attenuation Crosstalk Ratio*) : mesure, exprimée en décibels, du rapport signal/bruit d'un câble.

ADSL (*Asymmetric Digital Subscriber Line*) : technique de transmission à haut débit sur cuivre (de 1,5 à 8 Mbit/s dans un sens et 16 à 640 Kbit/s dans l'autre) sur une portée de 3,7 à 5,4 km. Permet de connecter les particuliers à l'Internet de façon permanente et à haut débit.

AMRC (*Accès multiples à répartition en code*) : désigne les techniques de multiplexage par attribution d'un code d'identification. Utilisé dans les réseaux cellulaires UMTS.

AMRF (*Accès multiples à répartition de fréquence*) : désigne les techniques de multiplexage fréquentiel. Utilisé dans les réseaux cellulaires analogiques.

AMRT (*Accès multiples à répartition dans le temps*) : désigne les techniques de multiplexage temporel. Utilisé dans les réseaux cellulaires numériques de type GSM.

APPN (*Advanced Peer to Peer Network*) : évolution des réseaux SNA vers un réseau non hiérarchique.

Arcnet (*Attached Ressource Computer Network*) : réseau local de type bus à jeton sur câble coaxial en étoile. Le débit est de 2,5 Mbit/s.

ARP (*Address Resolution Protocol*) : protocole basé sur un *broadcast* permettant d'obtenir l'adresse MAC (niveau 2) à partir d'une adresse réseau (niveau 3).

Arpanet (*Advanced Research Projects Agency Network*) : historiquement le premier réseau de type TCP/IP développé pour le DoD (*Department Of Defense*) américain.

ART (*Autorité de régulation des télécommunications*) : organisme de régulation dans le domaine des télécommunications en France (tarifs, concurrence, obligations des opérateurs, etc.).

ASIC (*Application-Specific Integrated Circuit*) : circuit intégré développé spécifiquement pour une application donnée et à la demande. Il est constitué d'une matrice de transistors qui forment des circuits logiques (AND, OR, XOR...) agissant sur des signaux en entrée et générant des signaux en sortie.

ASCII (*American Standard Code for Information Interchange*) : code sur 7 ou 8 bits, utilisé dans l'informatique pour représenter un caractère alphanumérique (128 ou 256 combinaisons sont possibles). Ce code est également normalisé ITU-T n° 5.

ASN.1 (*Abstract Syntax Notation 1*) : langage normalisé ISO permettant de décrire diverses structures comme les bases de données et les paquets des protocoles.

ATM (*Asynchronous Transfer Mode*) : protocole haut débit reposant sur la commutation de cellules de 53 octets.

AUI (*Attachment Unit Interface*) : type de connecteur à 15 broches utilisé pour les matériels Ethernet.

Autocom (abréviation de *autocommutateur*) : désigne un ordinateur spécialisé dans la commutation de circuits. Établit automatiquement une communication téléphonique en fonction d'un numéro de téléphone.

B

Backbone (*épine dorsale*) : désigne un réseau fédérateur à hauts débits permettant d'interconnecter des réseaux secondaires.

Balun (*Balanced Unbalanced*) : connecteur permettant d'adapter l'impédance entre deux câbles de nature différente.

Baud (du nom de *Émile Baudot*, inventeur du code télégraphique) : unité exprimant le nombre de modulations par seconde. Elle équivaut au bit par seconde si un signal représente une valeur binaire.

Bluetooth (du nom d'un roi norvégien) : technologie de réseau sans fil sur courte de distances dans le but connecter des périphériques à un ordinateur.

BBS (*Bulletin Board Systems*) : messagerie pour les PC permettant de télécharger des fichiers.

BERT (*Bit Error Rate Tester*) : test de la qualité d'une ligne consistant à générer des trames et à mesurer le taux d'erreur.

BISDN (*Broadband Integrated Services Digital Network*) : désigne le réseau numérique à haut débit reposant sur ATM et un support de transmission SDH (ou Sonet).

Bit (*Binary digiT*) : représente le plus petit élément d'information. Il prend les valeurs binaires « 1 » et « 0 ».

BNC (*Basic Network Connector*) : connecteur propre aux câbles coaxiaux utilisés par les réseaux Ethernet.

BOC (*Bell Operating Company*) : compagnies de téléphone américaines régionales au nombre de sept issues du démantèlement d'ATT (*American Telephone and Telecommunication*).

Bit/s (*bits par seconde*) : nombre de bits transmis par seconde.

BRI (*Basic Rate Interface*) : désigne l'accès de base RNIS (2B+D).

BSC (*Binary Synchronous Communications*) : protocole synchrone utilisé par certains équipements d'IBM.

BUS (*Broadcast and unknown Server*) : couche logicielle permettant de gérer les *broadcasts* émis par les réseaux locaux sur des chemins virtuels ATM.

C

CAP (*Carrierless Amplitude Phase*) : méthode de codage en ligne utilisée pour les réseaux hauts débits, les liaisons xDSL. Basé sur une quadruple modulation d'amplitude (QAM). Moins performant que DMT, mais moins cher.

CATV (*Cable Antenna TV*) : désigne la télévision par câble et tous les dispositifs s'y rattachant (câble, modulateur...).

CBDS (*Connectionless Broadband Data Service*) : transmission de données pour les réseaux de télécommunications hauts débits tels que ATM.

CDDI (*Copper Distributed Data Interface*) : version de FDDI sur des câbles cuivre à paires torsadées de catégorie 5.

CDMA (*Code Division Multiple Access*) : désigne les techniques de multiplexage par attribution d'un code d'identification. Utilisé dans les réseaux cellulaires UMTS. (cf. AMRC)

CICS (*Customer Information Control System*) : moniteur transactionnel d'IBM.

CISC (*Complex Instruction Set Component*) : type de microprocesseur offrant un grand nombre d'instructions (par opposition à RISC).

CMOS (*Complementary Metal Oxyde Semiconductor*) : technologie de fabrication de puces électroniques à base de silicium.

CPU (*Central Processor Unit*) : unité centrale exécutant les instructions d'un programme.

CRC (*Cyclic Redundancy Check*) : mécanisme de contrôle d'erreurs basé sur le calcul d'un polynôme générateur. Permet de contrôler que les données d'un paquet, d'une trame ou d'une cellule n'ont pas été endommagées lors de la transmission sur le réseau.

CSU (*Channel Service Unit*) : désigne un modem numérique raccordant l'équipement terminal (un routeur par exemple) à un nœud du réseau, comme un commutateur ATM (voir également DSU).

CSMA-CD (*Carrier Sense Multiple Access - Collision Detection*) : méthode d'accès au support physique par écoute du réseau et propre au réseau Ethernet.

CSMA-CR (*CSMA - Contention Resolution*) : méthode d'accès au support physique par gestion de priorité et propre à l'accès de base du RNIS.

CT2 (*Cordless Telephone 2nd generation*) : norme de radiotéléphones numériques sans fil. Exemple du Bi-Bop en France.

CV (*Circuit Virtuel*) : désigne un lien établi lors d'une procédure de connexion à travers plusieurs commutateurs d'un réseau. Les CVC (*circuits virtuels commutés*) sont établis à la demande. Les CVP (*circuits virtuels permanents*) sont établis une fois pour toutes lors de l'initialisation des équipements d'extrémité.

D

DECT (*Digital European Cordless Telecommunication*) : norme européenne de radiotéléphones numériques sans fil.

DES (*Data Encryption Standard*) : norme de cryptage spécifiée par l'ANSI (*American National Standards Institute*) et par le NIST (*National Institute of Standards and Technology*).

DHCP (*Dynamic Host Configuration Protocol*) : permet, à partir d'un serveur, de télécharger la configuration réseau vers un ordinateur (adresse IP, paramètres TCP/IP, etc.).

DNS (*Domain Name System*) : service de noms reposant sur des serveurs. Permet de convertir un nom en une adresse IP.

DLSw (*Data-Link Switching*) : désigne une méthode d'encapsulation des trames SNA dans des paquets TCP/IP (RFC 1434). Les acquittements des trames sont locaux.

DTMF (*Dual Tone MultiFrequency*) : système de signalisation utilisé pour la numérotation sur les réseaux téléphoniques analogiques.

DMT (*Discrete MultiTone*) : méthode de codage en ligne utilisée pour les réseaux hauts débits, les liaisons spécialisées et l'ADSL. Basé sur une quadruple modulation d'amplitude (QAM). Plus performant que CAP.

DMTF (*Desktop Management Task Force*) : groupe de travail ayant défini la norme du même nom qui a pour objet de décrire les protocoles et bases de données dans le domaine de l'administration bureautique.

DNA (*Digital Network Architecture*) : spécification des réseaux de la société DEC (*Digital Equipment Corporation*).

DQDB (*Distributed Queue Dual Bus*) : désigne la norme IEEE 802.6 pour les réseaux MAN (*Metropolitan Area Network*).

DSA (*Distributed System Architecture*) : modèle d'architecture réseau de la société Bull.

DSL (*Digital Subscriber Line*) : ensemble de technologies de transmission numérique à haut débit (quelques Mbit/s) fonctionnant sur des paires de cuivre comme celles du téléphone.

DSU (*Data Service Unit*) : désigne un adaptateur entre l'interface physique d'un ETTD et un support de transmission numérique tel qu'une LS (voir également CSU). Le protocole d'échange entre les deux est référencé sous le nom de DXI (*Data eXchange Interface*, RFC 1483).

E

EBCDIC (*Extended Binary Coded Decimal Interchange Code*) : code sur 8 bits permettant de coder les caractères alphanumériques.

Edge device (*Équipement de bordure*) : commutateur réseau local (Ethernet ou Token-Ring) disposant d'une interface ATM. Selon les cas, supporte les protocoles LANE, MPOA et PNNI. Dans certains cas, également appelé commutateur multiniveau.

ECMA (*European Computer Manufacturer Association*) : organisme de normalisation agissant dans les domaines des ordinateurs et des réseaux locaux.

EDI (*Electronic Data Interchange* ou *Échange de Données Informatisées*) : procédure d'échange de documents sous forme électronique. Spécifie la structure des données informatiques.

ELAN (*Emulated Local Area Network*) : réseau virtuel sur ATM.

EMC (*Electromagnetic Compatibility*) : mesure qui caractérise l'aptitude des différents composants informatiques, dont les câbles, à ne pas perturber d'autres composants.

ETEBAC (*Échanges télématiques entre les banques et leurs clients*) : ensemble de protocoles spécifiés par les organismes bancaires français pour l'échange de données informatiques.

ETSI (*European Telecommunication Standards Institute*) : organisme de normalisation participant aux travaux de l'ITU-T.

F

FAX (*Fac-similé*) : document télécopié.

FAI (*Fournisseur d'accès à Internet*) : désigne un opérateur réseau qui permet aux entreprises et aux particuliers de se connecter à l'Internet (*voir* ISP).

FCC (*Federal Communications Commission*) : organisme de régulation dans le domaine des télécommunications. Équivalent de l'ART en France.

FCS (*Frame Check Sequence*) : code de contrôle d'erreur équivalent au CRC.

FDDI (*Fiber Distributed Data Interface*) : réseau local haut débit (100 Mbit/s) en anneau à jeton spécifié par l'ANSI. FDDI-II est une version améliorée offrant le support de la voix et de la vidéo.

FEXT (*Far End Cross Talk*) : mesure, exprimée en décibels, du taux de réflexion dans un câble. Également appelé télédiaphonie.

Frame Relay (*Relais de trames*) : protocole couvrant les couches 2 et 3 du modèle OSI utilisé sur les liaisons longues distances. Considéré comme un protocole X25 allégé.

G

Gbit/s (*Gigabits par seconde*) : nombre de milliards de bits transmis par seconde.

GFA (*Groupe fermé d'abonnés*) : ensemble de nœuds interconnectés formant virtuellement un seul réseau inaccessible de l'extérieur bien que partageant la même infrastructure que d'autres clients au sein du réseau d'un opérateur.

Go (gigaoctet) : équivalent à 1 073 741 824 octets ($1\,024^3$).

GPRS (*General Packet Radio Service*) : évolution de la norme GSM permettant la transmission de données rapide.

GSM (*Global System for Mobile Communications*) : norme européenne définissant les réseaux téléphoniques numériques sans fil (réseaux cellulaires).

H

H.323 (*Protocole 323, Série H de l'ITU-T*) : architecture et protocole permettant d'établir des conversations téléphoniques sur un réseau IP (cf. VoIP et SIP).

HDB3 (*High-Density Bi-polar modulus 3*) : méthode de codage utilisée dans les liaisons MIC (avis ITU-T G703).

HDLC (*High-level Data Link Control*) : désigne une méthode d'encapsulation des données sur un support de transmission synchrone comme une ligne spécialisée.

HDSL (*High-bit-rate Digital Subscriber Line*) : technique de transmission à haut débit sur cuivre (1,544 et 2,048 Mbit/s sur une portée de 4,6 km) en *full duplex* à la différence de l'ADSL.

Hz (*Hertz*) : unité de fréquence correspondant à un cycle par seconde.

HPPI (*High Performance Parallel Interface*) : norme d'interface parallèle fonctionnant à 800 Mbit/s en transmettant 32 bits en parallèle ou fonctionnant à 2,6 Gbit/s sur 64 bits.

HTML (*HyperText Markup Language*) : langage de description d'une page web. Le navigateur (browser web) exécute ces commandes pour afficher la page.

HTTP (*HyperText Transfer Protocol*) : protocole permettant d'envoyer une page web d'un serveur web vers un ordinateur équipé d'un navigateur. Protocole à la base du web.

I

IETF (*Internet Engineering Task Force*) : organisme de standardisation des protocoles TCP/IP regroupant les constructeurs et opérateurs réseaux. Émet les RFC (*Request for Comments*).

INRIA (*Institut national de la recherche en informatique et en automatique*) : organisme français notamment en charge de la gestion de l'Internet en France.

Internet (*Inter Network*) : interconnexion de réseaux IP formant le plus grand réseau public à commutation de paquets du monde.

Intranet (*Intra Network*) : version privée de l'Internet dans une entreprise.

IP (*Internet Protocol*) : protocole de niveau 3 (couche réseau) à commutation de paquet. Protocole à la base de l'Internet.

IPX (*Internetwork Packet Exchange*) : protocole réseau utilisé par les serveurs Netware. IPX/SPX est l'équivalent de TCP/IP.

IPNS (*ISDN PABX Networking Specification*) : norme spécifiant la signalisation entre auto-commutateurs privés. Permet d'interconnecter des PABX de marques différentes.

ISDN (*Integrated Services Digital Network*) : désigne le réseau téléphonique numérique (voir RNIS).

IP-SEC (*IP Security*) : protocole IP sécurité. Les paquets sont cryptés et authentifiés. Standard à la base des VPN-IP.

ISP (*Internet Service Provider*) : désigne un opérateur réseau qui permet aux entreprises et aux particuliers de se connecter à l'Internet (cf. FAI).

ITSP (*Internet Telephony Service Provider*) : opérateur proposant des services VoIP (voix sur IP) sur Internet.

ITU-T (*International Telecommunication Union – Telecommunication standardization sector*) : organisme affilié à l'ONU qui définit les normes et la réglementation en matière de réseaux de télécommunications.

J

Jitter (*gigue*) : désigne le déphasage des signaux d'horloge dû à la distorsion des signaux sur la ligne.

JPEG (*Joint Picture Expert Group*) : groupe de travail et normes relatives à la compression d'images fixes.

K

Kbit/s (*Kilobits par seconde*) : nombre de milliers de bits transmis par seconde.

Ko (*Kilo-octet*) : équivalent à 1 024 octets.

KHz (*Kilohertz*) : désigne le nombre de milliers de cycles par seconde (fréquence) d'une onde ou d'une horloge.

L

LAN (*Local Area Network*) : désigne les techniques de réseau local.

LANE (*LAN Emulation Protocol*) : protocole permettant d'émuler un réseau local sur ATM. Utilise les composants LEC, LECS, LES et BUS.

LAT (*Local Area Transport*) : protocole, développé par Digital Equipment, permettant de relier des terminaux et imprimantes aux serveurs VAX en environnement Decnet.

LDAP (*Lightweight Directory Access Protocol*) : version allégée de l'annuaire X.500 et adaptée aux réseaux intranet et Internet.

LEC (*Lan Emulation Client*) : couche logicielle permettant d'adapter les protocoles de réseaux locaux sur ATM.

LECS (*Lan Emulation Configuration Server*) : couche logicielle contrôlant les ELAN. Gère les LES et indique aux *edge device* et stations ATM de quel LES ils dépendent en fonction de leur ELAN.

LES (*Lan Emulation Server*) : couche logicielle contrôlant un ELAN. Gère tous les *edge device* ou stations ATM appartenant à son ELAN ainsi que toutes les adresses MAC qui y sont rattachées côté réseau local.

LLC (*Logical Link Control*) : couche logicielle qui a pour objet d'assurer le transport des trames entre deux stations. Elle se situe au niveau 2 du modèle OSI et est fonctionnellement proche du protocole HDLC.

LS (*ligne spécialisée*) : ensemble de liaisons permanentes et vues du client comme étant point à point.

LU (*Logical Unit*) : dans les réseaux SNA, désigne une entité définissant des droits d'accès et des règles de communication avec d'autres entités. La LU 6.2 décrit par exemple les fonctions de communication entre les programmes.

M

MAC (*Medium Access Control*) : couche logicielle qui a pour rôle de structurer les bits d'information en trames adaptées au support physique et de gérer les adresses physiques des cartes réseaux (on parle d'adresse MAC).

MAN (*Metropolitan Area Network*) : désigne un réseau étendu, généralement en fibre optique, à l'échelle d'un campus ou d'une ville.

MAP (*Manufacturing Automation Protocol*) : réseau local spécifié par General Motors pour les environnements industriels. La méthode d'accès est celle du jeton sur un câble CATV. Proche de la norme IEEE 802.4.

Mbit/s (*Mégabits par seconde*) : nombre de millions de bits transmis par seconde.

MHS (*Message Handling System*) : sous-ensemble de la norme X400 spécifiant les mécanismes de gestion des messages.

MHz (*Mégahertz*) : désigne le nombre de millions de cycles par seconde (fréquence) d'une onde ou d'une horloge.

MIB (*Management Information Base*) : base de données structurée selon la syntaxe ASN.1 décrivant les objets d'un équipement réseau.

MIC (*Modulation par impulsions codées*) : technique de transmission utilisée pour véhiculer des signaux analogiques sous forme numérique par échantillonnage des signaux. Par extension, désigne ce type de ligne utilisé en France et comportant 32 canaux de 64 Kbit/s.

Modem (*modulateur-démodulateur*) : appareil transmettant des signaux analogiques sur le réseau téléphonique. Offre les fonctions de numérotation, de connexion, et éventuellement de compression et de correction d'erreur.

Mo (*Mégaoctet*) : équivalent à 1 048 576 octets ($1\,024^2$).

MPEG (*Moving Picture Expert Group*) : groupe de travail et normes relatives à la compression d'images animées.

MTA (*Message Transfer Agent*) : élément logiciel qui achemine les messages entre les différents nœuds d'un système de messagerie.

MTBF (*Mean Time Between Failure*) : mesure statistique, exprimée en nombre d'heures, du temps de fonctionnement d'un équipement avant une panne.

MTS (*Message Transfer System*) : désigne dans la norme X400 l'ensemble des MTA d'un même domaine.

MTTR (*Mean Time To Repair*) : exprime le temps de rétablissement suite à une panne.

MTU (*Maximum Transmission Unit*) : longueur (ou taille) maximale d'une trame ou d'un paquet d'un protocole réseau.

N

NDIS (*Network Driver Interface Specification*) : spécification par Microsoft d'une interface logicielle universelle d'accès aux cartes réseau (NIC) d'un PC. Permet au logiciel situé au niveau de la couche 2 du modèle OSI d'utiliser n'importe quelle carte réseau (voir ODI).

Netbios (*Network Basic Input/Output System*) : protocole de niveau session permettant de partager des ressources entre postes de travail et serveurs en environnement Windows (95, 98 ou NT).

NEXT (*Near-End Cross(X) Talk*) : mesure de la paradiaphonie d'un câble cuivre (aptitude d'un câble à ne pas être perturbé par les parasites).

NFS (*Network File System*) : système de gestion de fichiers réseau sur TCP/IP. Permet de partager des fichiers en donnant à l'utilisateur l'impression qu'ils sont locaux.

NIC (*Network Interface Card*) : désigne une carte réseau dans un ordinateur.

NOS (*Network Operating System*) : désigne les systèmes d'exploitation des serveurs de fichiers tels que Netware ou Windows NT.

NRZ (*Non Return to Zero*) : méthode de codage des signaux numériques.

NTI (*Nœud de Transit International*) : commutateur assurant l'interconnexion des réseaux X25 entre les différents pays.

O

ODI (*Open Data-link Interface*) : spécification par Novell d'une interface logicielle universelle d'accès aux cartes réseau (NIC) d'un PC. Permet au logiciel situé au niveau de la couche 2 du modèle OSI d'utiliser n'importe quelle carte réseau (voir NDIS).

OEM (*Original Equipment Manufacturer*) : contrat de fabrication sous licence.

OLE (*Object Linking and Embedding*) : protocole d'origine Microsoft pour la communication des applications réparties.

OLTP (*On Line Transaction Processing*) : désigne un moniteur transactionnel en ligne.

OSPF (*Open Shortest PAth First*) : protocole de routage TCP/IP.

P

PABX (*Private Automatic Branch eXchange*) : appelé autocommutateur ou autocom, il assure la concentration des postes téléphoniques et la commutation des circuits.

PAD (*Paquet Assembler / Desassembler*) : équipement permettant aux terminaux asynchrones (travaillant caractère par caractère) d'accéder à un réseau de paquets X25.

PAV (*Point d'accès Vidéotex*) : variante du PAD pour l'accès aux services Minitel reposant sur le réseau X25 de Transpac.

PC (*Personal Computer*) : micro-ordinateur. Désigne un ordinateur compatible avec les spécifications des sociétés IBM et Intel en matière d'ordinateur personnel.

PCM (*Pulse Code Modulation*) : équivalent américain du MIC (Modulations par Impulsions Codées) comportant 24 canaux à 64 Kbit/s.

PESIT (*Protocole d'Échange pour le Système Interbancaire de Télécompensation*) : ensemble de procédures spécifiées par les banques françaises pour les échanges de données électroniques entre les banques.

Ping (*Packet Internet Groper*) : paquet de la couche réseau utilisé pour mesurer les temps de réponse d'un réseau. L'émetteur attend le même paquet en écho.

Pixel (*Picture Element*) : désigne un point élémentaire dans une image.

PNNI (*Private to Private Network*) : protocole de routage ATM. Permet aux commutateurs ATM de déterminer les meilleurs chemins virtuels.

PPP (*Point-to-Point Protocol*) : protocole de la couche liaison utilisé sur les lignes série téléphoniques ou spécialisées.

PRI (*Primary Rate Interface*) : désigne l'accès primaire RNIS.

PSDN (*Packet Switching Data Network*) : désigne un réseau à commutation de paquets.

PSTN (*Public Switched Telephone Network*) : désigne le réseau téléphonique. Équivalent au RTC français (réseau téléphonique commuté).

Q

Q-SIG (*Q signalisation*) : norme basée sur la signalisation CITT Q.931 définissant les échanges entre les systèmes de signalisation publics et privés (PABX).

R

RADSL (*Rate-adaptive Asymmetric Digital Subscriber Line*) : technique de transmission à haut débit analogue à ADSL mais avec modification du débit en fonction de la qualité de la ligne.

RARP (*Reverse Address Resolution Protocol*) : protocole basé sur un *broadcast* permettant d'obtenir l'adresse réseau (niveau 3) à partir d'une adresse MAC (niveau 2).

RFC (*Request For Comments*) : documents issus de l'IETF (*Internet Engineering Task Force*) spécifiant tous les standards en matière de protocoles Internet.

RISC (*Reduced Instruction Set Components*) : type de microprocesseur caractérisé par un jeu d'instructions réduit au minimum (opposé à CISC).

RIP (*Routing Information Protocol*) : protocole de routage TCP/IP et IPX/SPX.

RNIS (*réseau numérique à intégration de services*) : désigne le réseau téléphonique numérique censé remplacer le RTC petit à petit (cf. ISDN).

RPIS (*réseau privé à intégration de services*) : interconnexion de PABX et de lignes pour former un réseau RNIS privé.

RPV (*réseau privé virtuel*) : voir VPN.

RSA (*Rivest Shamir et Adelman*) : initiales des inventeurs de l'algorithme de chiffrement du même nom.

RSVP (*Resource Reservation Protocol*) : protocole de signalisation permettant de garantir une qualité de service sur les réseaux TCP/IP pour les applications temps réel comme le transport de la voix.

RTC (*réseau téléphonique commuté*) : désigne le réseau téléphonique analogique classique.

RVA (*réseau à valeur ajoutée*) : désigne un réseau associé à des services tels que la conversion de protocoles, la facturation, l'accès à des bases de données, etc.

S

SAP (*Service Access Point*) : mécanisme logiciel de pointeurs permettant à un logiciel réseau d'utiliser les services d'une couche inférieure. Le SAP est un numéro unique permettant d'identifier le logiciel qui a envoyé une trame ou un paquet.

SAN (*Storage Area Network*) : réseau haut débit sur fibre optique, reposant sur la technologie Fibre Channel, et dédié aux sauvegarde des données.

SDH (*Synchronous Digital Hierarchy*) : mode de transmission numérique pour les réseaux de télécommunications hauts débits.

SDSL (*Single-line Digital Subscriber Line*) : technique de transmission à haut débit sur cuivre (1,544 et 2,048 Mbit/s sur une portée de 3 km) similaire à HDSL.

SIP (*Session Initiation Protocol*) : architecture et protocole de l'IETF permettant d'établir des conversations téléphoniques sur un réseau IP (cf. VoIP et H.323).

SIT (*Système Interbancaire de Télécompensation*) : réseau d'échange de données entre les banques françaises.

SMDS (*Switched Multimegabit Data Service*) : réseau à commutation de paquets à hauts débits pour des liaisons longues distances.

SMTP (*Simple Mail Transfert Protocol*) : protocole utilisé par les systèmes de messagerie dans le monde TCP/IP.

SNA (*Systems Network Architecture*) : spécifications des réseaux de la société IBM (*International Business Machine*).

SNMP (*Simple Network Management Protocol*) : protocole de la famille TCP/IP utilisé pour administrer à distance les équipements réseaux à partir d'une station d'administration.

SPX (*Sequenced Packet Exchange*) : protocole réseau utilisé par les serveurs Netware. IPX/SPX est l'équivalent de TCP/IP.

STP (*Shielded Twisted Pair*) : désigne un câble blindé composé de 2 ou 4 paires en cuivre (voir UTP).

T

TASI (*Time Assignment Speech Interpolation*) : technique de multiplexage temporel statistique adaptée à la transmission de la voix numérisée et utilisée dans les satellites et les câbles sous-marins.

TAXI (*Transparent Asynchronous Transmitter/Receiver Interface*) : spécification d'une interface fibre optique à 100 Mbit/s pour FDDI et ATM. Utilise le codage 4B/5B.

TCP/IP (*Transport Control Protocol / Internetwork Protocol*) : protocole de transport des données sous forme de paquets, universellement utilisé sur les réseaux LAN et WAN. Désigne également toute une famille de protocoles de niveau session ou application (HTTP, FTP, SMTP, SNMP, etc.).

TDM (*Time Division Multiplexing*) : technique de multiplexage dans le temps (voir AMRT).

TIC (*Token-Ring Interface Coupler*) : dénomination d'un mode d'attachement des contrôleurs IBM SNA 3174, 3745, etc. Désigne une connexion SNA sur réseau local.

TNR (*Terminaison Numérique de Réseau*) : coffret marquant la séparation entre le réseau RNIS public et la partie privée chez l'utilisateur. Correspond à la prise téléphonique.

TOP (*Technical Office Protocol*) : ensemble de protocoles développés par la société Boeing pour la CAO (*conception assistée par ordinateur*). TOP respecte entièrement la normalisation OSI.

TTL (*Time To Live*) : compteur permettant de déterminer le temps de validité restant pour une donnée ou un paquet réseau.

U

UA (*User Agent*) : partie cliente des systèmes de messagerie.

UHF (*Ultra High Frequency*) : bande de fréquence, située entre 30 MHz et 300 MHz, utilisée pour transmettre des émissions de télévision analogiques.

UDP (*Network Interface Card*) : équivalent de TCP mais en mode non connecté, sans les mécanismes de contrôle de flux, de reprise sur erreur et autres options.

UTP (*Unshielded Twisted Pair*) : désigne un câble non blindé composé de 4 paires en cuivre. La norme EIA-TIA 586 définit 5 catégories de câbles de ce type (voir STP).

V

VCC (*Virtual Channel Connection*) : chemins virtuels ATM prédéfinis ou affectés dynamiquement pour les besoins des protocoles LANE et MPOA. Un VCC héberge une encapsulation LLC telle que définie par la RFC 1483.

VDSL (*Very-high-bit Digital Subscriber Line*) : technique de transmission à haut débit sur cuivre (de 13 à 52 Mbit/s sur une portée de 300 à 1 300 mètres) similaire à ADSL.

VHF (*Very High Frequency*) : bande de fréquence, située entre 300 MHz et 3 GHz, utilisé pour transmettre des émissions de télévision analogiques.

VLAN (*Virtual Local Area Network*) : désigne les réseaux locaux virtuels. L'extension des VLAN sur ATM passe par les ELAN.

VoIP (*Voice Over IP*) : désigne les technologies permettant de transmettre des conversations téléphoniques (et des visioconférences) sur un réseau IP. Repose sur les protocoles H.323 et SIP.

VPN (*Virtual Private Network*) : désigne un réseau WAN (reposant sur ATM, Frame-Relay) dédié à une entreprise mais reposant sur un backbone haut débit que l'opérateur partage avec d'autres clients. L'opérateur garantit que les données ne sont pas mélangées avec celles d'un autre client.

VPN-IP (*Virtual Private Network – Internet Protocol*) : liaison IP protégée par cryptage des données.

VSAT (*Very Small Aperture Terminals*) : protocole de transmission satellite utilisant des paraboles de faibles diamètres (inférieurs à 3,7 m).

W

WAN (*Wide Area Network*) : désigne un réseau longue distance (réseau étendu) reposant sur les technologies de transmissions de données en série et des protocoles tel que le Frame Relay et ATM.

WDM (*Wavelength Division Multiplexing*) : mode de transmission numérique sur fibre optique multiplexant différentes longueurs d'onde et autorisant de très hauts débits (> 100 Gbit/s).

Web (*la toile*) : désigne l'ensemble des serveurs de page web de l'Internet.

WLAN (*Wireless LAN*) : réseau local sans fil à 11 Mbit/s, compatible Ethernet, et normalisé IEEE 802.11.

WML (*Wireless Markup Language*) : protocole de description d'une page web (version simplifiée de HTML et de HTTP) et adapté aux faibles débits des téléphones GSM.

WWW (*World Wide Web*) : désigne l'ensemble des serveurs web de l'Internet. Nom DNS généralement donné aux serveurs web.

X

X25 (*norme X.25 de l'ITU-T*) : protocole de niveau 3 spécifiant l'interface d'accès à un réseau à commutation de paquets.

X400 (*norme X.400 de l'ITU-T*) : systèmes de messagerie OSI (défini les protocoles et les logiciels MTA, UA)

X500 (*norme X.500 de l'ITU-T*) : annuaire de messagerie OSI.

XNS (*Xerox Network System*) : l'un des tous premiers protocoles LAN de la société Xerox, équivalent de TCP/IP et IPX/SPX.

xDSL (*x Digital Subscriber Line*) : regroupe les techniques de transmission ADSL, RADSL, HDSL, SDSL et VDSL.

Bibliographie

- U. BLACK, *OSI – A Model for Computer Communications Standards*, Prentice Hall, 1991.
- M. BOISSEAU, M. DEMANGE, J.-M. MUNIER, *Réseaux ATM*, Eyrolles, 2^e édition, 1996.
- CCITT, Livre bleu, 1998.
- D. COMER, *Internetworking with TCP/IP – Principles, Protocols and Architecture*, Prentice Hall, 1991 (trad. fr. : *TCP/IP – Architecture, protocoles, applications*, InterÉditions, 1992).
- M. DE PRYCKER, *Asynchronous Transfer Mode – Solution for Broadband ISDN*, Ellis Horwood, 1993 (trad. fr. : *ATM – Mode de transfert asynchrone*, Masson, 1994).
- ÉLECTRICITÉ DE FRANCE, *Principes de conception et de réalisation des mises à la terre*, H 115, janvier 1984.
- ÉLECTRICITÉ DE FRANCE, *Pour avoir un bon réseau de terre : des règles simples pour respecter la physique de base*, RGE n° 10, novembre 1986.
- A. FENYÖ, F. Le GUERN, S. TARDIEU, *Raccorder son réseau d'entreprise à l'Internet*, Eyrolles, 1997.
- FRANCE TÉLÉCOM, *STUM : spécifications techniques d'utilisation du Minitel*, 1990.
- FRANCE TÉLÉCOM, *STUR : spécifications techniques d'utilisation du Réseau*, 1990.
- FRANCE TÉLÉCOM, *STAS : spécifications techniques d'accès aux services Numéris*, 1990.
- C. HUITEMA, *Le routage dans l'Internet*, Eyrolles, 1995.
- C. MACCHI, J.-F. GUILBERT, *Téléinformatique*, Dunod, 1988.
- J-L. MÉLIN, *Pratique des réseaux ATM*, Eyrolles, 1997
- J-L. MONTAGNIER, *Pratique des réseaux d'entreprise*, Eyrolles, 1999
- G. PUJOLLE, *Les Réseaux*, Eyrolles, 1995.
- P. ROLIN, *Réseaux locaux : normes et protocoles*, Hermès, 1988.
- G.C. SACKETT, C.Y. METZ, *ATM and Multiprotocol Networking*, McGraw-Hill, 1997.
- W. STALLINGS, *Handbook of Computer-Communications Standards – Volume 1 : The OSI Model*, Macmillan, 1990.
- W. STALLINGS, *Handbook of Computer-Communications Standards – Volume 2 : LAN Standards*, Macmillan, 1990.
- W. STALLINGS, *Handbook of Computer-Communications Standards – Volume 3 : The TCP/IP Protocol Suite*, Macmillan, 1990.
- A. TANENBAUM, *Computer networks*, Prentice Hall, 1988 (trad. fr. : *Réseaux – Architectures, protocoles, applications*, InterÉditions, 1990).

Câblage

www.eia.org	☐ <i>Electronic Industries Alliance</i>
www.tiaonline.org	☐ <i>Telecommunications Industry Association</i>
www.broadband-guide.com	
www.scope.com/standards	☐ <i>Testeurs de câblage Wirescope</i>
www.cablingstandards.com	
www.necanet.org	☐ <i>National Electrical Contractors Association</i>
www.nema.org	☐ <i>National Electrical Manufacturers Association</i>
www.rs232.net	☐ <i>Portail spécialisé sur le câblage</i>

Internet

www.afnic.asso.fr	☐ <i>Association française pour le nommage Internet en coopération</i>
www.apnic.net	☐ <i>Asia Pacific Network Information Centre</i>
www.arin.net	☐ <i>American Registry for Internet Numbers</i>
www.darpa.mil	☐ <i>Defence Advanced Research Projects Agency</i>
www.iab.org	☐ <i>Internet Architecture Board</i>
www.iana.org	☐ <i>Internet Assigned Numbers Authority</i>
www.icann.org	☐ <i>Internet Corporation for Assigned Names and Numbers</i>
www.ietf.org	☐ <i>Internet Engineering Task Force</i>
www.isoc.org	☐ <i>Internet Society : portail d'accès à tous les organismes régulant l'Internet</i>
www.internic.net	☐ <i>Enregistrement des noms de domaines</i>
www.register.com	
www.ripe.net	☐ <i>Réseau IP européen</i>
www.rfc-editor.org	☐ <i>Accès à toutes les RFC (Request for Comments)</i>
www.6bone.net	☐ <i>Réseau Internet IPv6</i>

www.mbone.com	☐ <i>Réseau Internet multicast</i>
www.renater.fr	☐ <i>Réseau national de la recherche</i>
www.internet2.edu	☐ <i>L'Internet dédié aux universités</i>
www.ucaid.edu/abilene	☐ <i>Abilene : l'Internet dédié aux universités</i>
www.cybergeography.org/mapping.html	☐ <i>Cartographie de l'Internet</i>
www.isc.org/ds	☐ <i>Internet Survey</i>

Modem-câble

www.opencable.com	
www.ietf.org/html.charters/ipcdn-charter.html	☐ <i>IP over Cable Data Network</i>
www.cabledatcomnews.com	☐ <i>Portail dédié aux modems câble</i>

Organismes de normalisation

web.ansi.org	☐ <i>American National Standards Institute</i>
www.etsi.org	☐ <i>European Telecommunications Standards Institute</i>
www.ema.org	☐ <i>Electronic Messaging Association</i>
www.iso.ch	☐ <i>International Standard Organization</i>
www.itu.ch	☐ <i>International Telecommunications Union</i>
www.ietf.org	☐ <i>Internet Engineering Task Force</i>
www.regulate.org	☐ <i>World Regulatory Telecommunications</i>
standards.ieee.org	☐ <i>Institute of Electrical and Electronics Engineers</i>
www.w3.org	☐ <i>WWW Consortium</i>

Protocoles

www.protocols.com	☐ <i>Description de tous les protocoles</i>
www.atmforum.com	☐ <i>ATM Forum</i>
www.frforum.com	☐ <i>Frame-Relay Forum</i>
www.gigabit-ethernet.org	☐ <i>Gigabit Ethernet Alliance</i>
www.fibrechannel.com	☐ <i>Fibre Channel Forum</i>

www.tl.org	☐ <i>Committee T1</i>
www.adsl.com	☐ <i>ADSL Forum</i>
www.ietf.org/html.charters/wg-dir.html	☐ <i>Portail d'accès à tous les groupes de travail produisant les RFC</i>
www.nmf.org	☐ <i>Network Management Forum</i>

Qualité de service

www.itmcenter.com	☐ <i>Internet Traffic management ressource Center</i>
diffserv.lcs.mit.edu	☐ <i>Portail d'accès aux information sur Diffserv</i>
www.cis.ohio-state.edu/~jain/refs/ipqs_ref.htm	
www.qosforum.com	☐ <i>Quality Of Service Forum</i>
http://www.ietf.org/html.charters/diffserv-charter.html	

Réseaux sans fils

www.mobinet.com	☐ <i>L'Internet mobile</i>
www.bluetooth.com	☐ <i>Portail spécialisé sur Bluetooth</i>
www.dectweb.com	☐ <i>Wireless Telecommunications Internet Services</i>
www.dect.ch	☐ <i>DECT Forum</i>
www.wlana.com	☐ <i>Wireless LAN Alliance</i>
www.umts-forum.org	☐ <i>UMTS Forum</i>

Revue de presse

www.reseaux-telecoms.fr	
www.01-informatique.com	
www.lmi.fr	☐ <i>Le Monde informatique</i>
www.data.com	☐ <i>Data Communications</i>
www.networkcomputing.com	
www.telecoms-mag.com	
www.cmpnet.com	☐ <i>Portail d'accès à de nombreuses revues spécialisées</i>
www.zdnet.com	☐ <i>Toute l'actualité informatique</i>

www.techguide.com	☐ <i>Portail d'accès aux technologies</i>
www.techweb.com	☐ <i>Portail d'accès aux technologies</i>
www.entmag.com	☐ <i>Windows NT & 2000 News</i>
www.lantimes.com	☐ <i>LAN times</i>
www.samag.com	☐ <i>SysAdmin : le journal des administrateurs Unix</i>

VoIP

www.ectf.org	☐ <i>Enterprise Computer Telephony Forum</i>
www.imtc.org	☐ <i>International Multimedia Telecommunications Consortium</i>
itel.mit.edu	☐ <i>Internet & Telecoms Convergence Consortium</i>
standard.pictel.com	☐ <i>Picturetel : accès aux standards H.323</i>
www.telephonyworld.com/iptelep/iptelep.htm	☐ <i>Portail VoIP</i>
www.theipsite.com	
www.von.com	☐ <i>Voice On The Net</i>
www.gvcnet.com	☐ <i>Portail spécialisé sur la vidéoconférence</i>
www.openh323.org	☐ <i>Informations sur les protocoles H.323</i>
www.mpeg.org	☐ <i>Moving Picture Experts Group</i>

Index

1

10bT, 100bT, 87, 96, 227

2

2B1Q (codage), 170, 174

8

802.1q, 228, 233, 302

A

AAL (*ATM Adaptation Layer*), 210, 251, 319

Administration réseau, 356

ADPCM (codage), 249

Adressage, 122

Adressage

ATM, 215, 222

Frame-Relay, 204, 223

IP, 42, 61, 111, 134, 151, 259

MAC, 111, 132, 133, 261, 366

plan, 365

public/privé, 113, 119

unicast/multicast, 259

ADSL, 8, 142, 173, 176

Affaiblissement/Atténuation, 71, 140

Agent SNMP, 356

Agrégation de canaux B, 158

Agrégation de liens Ethernet, 104, 230

Aire OSPF, 237, 274

A-law, μ -law, 249

Alias DNS, 390

AMI (codage), 170

Analogique, 141

Analyseur réseau, 354

Annuaire, 39, 181

ARP, 135, 155

ATM (*Asynchronous Transfer Mode*), 4, 41, 86, 97, 121, 141, 174, 189, 208, 224, 319, 407, 415

AUI, 52

Autocom, 25

Autosense, 87, 91

AWG, 174

B

Backbone, 97, 169

Baie (câblage), 79

BECCN (*Backward Explicit Congestion Notification*), 205

BERT (*Bit Error Tests*), 191

BGP, 274

BNC, 51, 78

Bootp, 365

Boucle locale, 165, 171, 175

Brassage, 53, 79, 147

Broadcast

IP, 111, 259

MAC, 133, 141, 155, 226, 370

Browser, 17, 26

de MIB, 362

Burst, 193, 292

C

Câble blindé/écrané, 72

Câble cuivre, 50

Cache

ARP, 135

DNS, 381

navigateur, 30, 182

Cahier des charges (câblage), 80

CAP (codage), 175

Carte réseau, 51, 57, 92, 133, 151, 363, 366

Catégorie 5, 6, 7, 74

CCTP, 66, 81

CDSL, 174

Cellule ATM, 174

CELP (codage), 249
 CEM, 74
 CERT, 45
 CHAP, 162
 CIF (*Common Intermediate Format*), 252
 CIR Frame Relay, 193, 205
 Circuit virtuel
 ATM, 212
 Frame-Relay, 199
 Classe D, E, F, 74
 Classe de service, 304, 312
 Classical IP, 190, 223, 225
 Classification, 301
 CLLM (*Consolidated Link Layer Management*), 205
 Coaxial (câble), 72, 76
 Codage, 170, 174, 248, 341
 Codec
 audio, 249, 254, 319
 vidéo, 250, 319
 Collapse backbone, 97, 224
 Collision, 90
 Commutateur
 niveau 2, 125
 niveau 3, 226, 233
 Commutation
 de cellule, 211, 318
 de circuit, 221, 253, 255, 318
 de paquet, 255, 318
 Ethernet, 125
 matrice, 92, 125
 Compilation de MIB, 360
 Compression, 248, 251
 Frame-Relay, 207
 RTP, 344
 TCP/IP, 157
 Concentrateur, 54, 89, 125
 Contrôle de flux
 Gigabit, 230
 Cookies, 32
 COPS, 314
 COS, 298
 Couches réseaux, 121, 125
 CPE, 165, 171
 CS-ACELP (codage), 249
 CSU (*Channel Service Unit*), 191
 CSU/DSU, 141, 171, 191, 208, 212, 213
 Cuivre/fibre optique (choix), 71, 101, 104, 227

D

DB9/DB25, 14, 17, 150
 DCE/DTE, 141, 150
 DCT (codage), 251
Dead gateway, 234
 Débit, 3, 13, 21, 87, 98, 143, 184, 193
 Débordement RNIS, 163
 DECT, 67, 335
Default gateway, 154, 233
 Délai de propagation, 140
 Délai de transit, 210, 220, 252, 312, 344
 DHCP, 157, 363
 Diffserv, 294
 DLCI Frame-Relay, 194, 202
 DMT (codage), 175
 DNS (*Domain Name System*), 27, 42, 329, 376, 381
 client, 395
 Driver, 16, 18, 60
 DS0 (canal 64 Kbit/s), 170, 248
 DSCP, 300, 302
 DSL (*Digital Subscriber Line*), 8, 141, 171, 175
DTMF, 253
 DVMRP, 265, 280
 DWMT (codage), 175

E

E.164, 223, 325
 E1/T1, 169, 170, 249
 Échantillonnage, 248, 341
 Echo, 252
 EIA/TIA, 402
 ELMI, 205
 E-mail, 34
 Encapsulation, 121, 254
 Ethernet, 136
 IP, 133
 Netbios, 138
 Erlang, 255
 Espace d'adressage, 122
 Ethernet, 71, 86, 92, 121, 125, 140
 Étoile, 50

F

FDDI, 97
 Fédérateur (réseau), 97, 231
 Fibre optique, 41, 71, 74, 84, 105, 174, 231
 Fichier DNS
 cache, 384
 initialisation, 394
 Fichier MIB, 359
 FIFO, 290
 File d'attente, 288, 293
 Firewall, 117
 FIRST, 45
 Flux
 client-serveur, 180
 conversationnel, 178
 transactionnel, 178
 FRAD (*Frame-Relay Access Device*), 191, 198
 Frame-relay, 41, 141, 188, 191
 FTP, 33, 288, 393
 Full duplex (Ethernet), 91

G

G.711, G.722, G.723, etc., 249, 254, 319
 Gatekeeper, 320, 329, 334
 Générateur de trafic, 355
 Gigabit, 74, 76, 96, 189, 230, 231
 Gigue, 220, 253, 314
 Groupe multicast, 261
 GTB (gestion technique du bâtiment), 67

H

H.225, 340
 H.245, 322, 340
 H.261, H.263, 251
 H.323, 78, 320
 HDLC, 152
 HDSL, 169, 172
 Hosts (fichier), 376
 HSRP (*Hot Standby Router Protocol*), 233
 HTML (*HyperText Markup Language*), 29
 HTTP, 124, 298
 Hub, 54, 89, 125

I

IAB (*Internet Architecture Board*), 43
 IANA (*Internet Assigned Numbers Authority*), 45, 259
 ICANN (*Internet Corporation for Assigned Names and Numbers*), 44
 ICMP, 289, 350, 351
 IDSL, 171
 IEEE, 404
 IEPG, 45
 IETF (*Internet Engineering Task Force*), 43
 IGMP (*Internet Group Membership Protocol*), 261, 266, 271, 276
 ILMI (*Integrated Local Management Interface*), 218
 Impédance, 72
 INRIA, 45
 Interface routeur, 147, 151, 161, 191, 202, 207
 Interface série, 151, 152, 196
 Interframe/Intraframe (compression vidéo), 251
 Internet, 12, 15, 27, 40, 42
 Intersites (réseau), 140, 176, 190, 192, 284
 Intranet, 390
 Intserv, 294, 305
 Inverse ARP, 214, 223
 IP precedence, 294, 301, 339
 Ipconfig, 372
 IPng, 409
 IPv6, 7, 43, 409, 436
 IRTF (*Internet Research Task Force*), 43
 ISOC (*Internet Society*), 43
 ISP (*Internet Service Provider*), 12, 15, 20, 27, 35, 41
 ITU, 164
 ITU-T (*International Telecommunication Union*) (CCITT), 43, 402

J

Jitter, 220, 252
 JPEG, 251

K

Kbit/s, 141

L

LAN, 8, 116, 236, 303
 LANE (*LAN Emulation*), 189, 225
 LAP-F, 199
 LDAP, 39
 LD-CELP (codage), 249
 Lien virtuel OSPF, 241
 LLC, 123
 LMI (*Local Management Interface*), 199, 205
 Longueur d'onde, 84
 LS (ligne spécialisée), 141, 158, 165, 169, 188, 190
 LTE (locaux techniques d'étagé), 68, 71, 78, 79, 148

M

MAC, 126, 131, 155, 234
 MAE (Internet), 41
 MAN, 8, 231
 Marquage, 300
 MBONE, 281, 287
 MCU H.323, 320
 Mean Option Score, 250
 Messagerie, 34, 389
 MIC (modulation par impulsions codées), 249, 417
 Mixer RTP, 345
 M-JPEG (Motion JPEG), 251
 MLQ (codage), 249
 Modem, 13, 18, 22, 25
 MOSPF, 270, 280
 MPEG, 251
 MPOA (*Multi Protocol Over ATM*), 225
 MTA (*Message Transfer Agent*), 389
 MTU (*Maximum Transfer Unit*), 313
 Multicast
 IP, 111, 259, 285, 307
 MAC, 141
 Multimédia, 3, 248, 340, 346
 Multimode/monomode, 75
 Multipaire (câble), 72
 Multiplexage
 des circuits virtuels ATM, 212
 des protocoles, 123
 temporel, 174
 Multiplexeur, 165, 169, 170
 MX (enregistrement DNS), 389

N

NAP (Internet), 41
 Navigateur, 17, 26, 33, 298
 Netbios, 58, 137, 138
 Netstat, 373
 NIC (*Network Information Center*), 46
 NLPID (IP dans Frame-Relay, 190, 223
 Nœud réseau, 22
 Nom DNS complet/relatif, 395
 Nommage, 42, 376
 Noms de domaine, 42
 Normalisation, 43
 NSAP, 218
 NSF (*National Science Foundation*), 45
 Nslookup, 394, 396
 Numérique, 141
 Numérisation, 248

O

Ohm, 72
 Opérateur, 12, 40, 143, 165, 191, 198
 OSPF (*Open Shortest Path First*), 237, 270, 274
 Outsourcing, 169
 Overhead, 141, 184, 189, 344

P

PABX, 67, 72, 197, 324, 327
 IP, 335
 Paire torsadée, 50, 76
 PAM (codage), 175
 Paquet
 IP, 123, 133, 142, 294
 TCP, 123, 178
 Paradiaphonie, 74
 Passerelle H.323, 320, 324, 334
 PCM (Pulse Code Modulation), 249
 PIM, 275, 280
 Ping, 157, 350
 Plan d'adressage, 110
 Plan de nommage, 376
 Plate-forme d'administration, 356
 Plésiochrone, 170
 Policing, 296, 297, 300
 Pont, 155
 POP (opérateur), 15, 165, 192, 198

Port (tcp/udp), 45, 123, 137, 301
 POTS (*Plain Old Telephone Service*), 326
 PPP, 21, 121, 142, 152, 162, 408, 411
 PPP (Internet), 41
 Priorité, 288
 Prise téléphonique, 14
 Prises RJ11/RJ45, 18
 Proxy
 ARP, 156
 cache Web, 31
 PVC
 ATM, 211, 215
 Frame-Relay, 194, 202

Q

Q.931, 322, 340
 QAM (codage), 175
 QCIF, 252
 QoS (Quality of Service), 287, 307
 Qualité de service, 3, 41, 167, 205, 208, 219, 255, 294, 312, 339
 Quantification, 248

R

RADSL, 173
 RAS (signalisation), 322, 340
 Recette (câblage), 66, 83
 Redondance, 103, 158, 233, 239
 Réflectométrie, 84
 Registry Windows
 dead gateway, 234
 default gateway, 368
 DNS, 383, 394
 TTL, 350
 Relais de trames, 416
 Relais DHCP, 370
 Résolution d'adresse, 135
 Résolution DNS inverse, 393
 Resolver DNS, 376, 395
 RFC (Request For Comments), 42
 RFC1700, 45, 123, 259
 RFC1918, 113, 119
 RIP, 265
 RJ11, 18
 RJ45, 18, 50, 53, 71, 72, 78, 94, 148, 150
 RLE (compression), 251

RNIS (réseau numérique à intégration de services), 141, 147, 171, 319, 413
 secours, 158
 Routage, 125, 152, 226, 233, 237
 Route (commande windows), 373
 Routeur, 142, 147, 151, 196, 202, 226, 234, 262, 296, 370
 configuration, 150
 RS-232, 150
 RSVP, 307, 313
 RTC, 14, 121, 141, 165, 253, 319
 RTP/RTCP, 340, 341

S

SC (connecteur), 71, 75, 78
 SDH (*Synchronous Data Hierarchy*), 141, 170, 174, 410, 417
 SDLC, 123
 SDSL/SHDSL, 173
 Secours RNIS, 158
 Segmentation, 56, 89
 Sérialisation, 253
 Série
 câble/port, 14, 17, 19, 148, 150
 interface, 148, 191, 208
 Serveur DNS, 381
 cache, 392
 primaire, 385
 racine, 390
 secondaire, 392
 Service opérateur, 145, 165
 Shannon, 248
 Signalisation, 170, 305, 319
 SIP, 7, 318
 SLA (*Service Level Agreement*), 167
 SLC (codage), 175
 SMTP, 389
 SNA, 123
 SNAP, 123, 190
 SNMP, 92, 119, 356, 408, 410
 Socket, 136
 Sonet, 170
 Spaming, 39
 Spanning tree, 104, 125, 128
 Spécificateur, 305
 Stackable, 56, 87
 Station d'administration, 356
 STM, 174

STP (câble), 74
STUN, 123
Subnet, 114, 116, 119, 120, 149, 151, 156, 233
Support de transmission, 144
SVC, 199
SVC
 ATM, 211, 214
 Frame-Relay, 194, 204
Switch Ethernet, 129
Synchrone, 170, 176

T

T.120, 319, 340
T1/E1, 169, 170, 248
Tableau virtuel, 346
TCP, 136
TCP/IP, 16, 57, 405, 406
 configuration, 22, 154, 363, 371
TDM, 174
Telnet, 133, 137, 144, 178, 288, 311
Temps de réponse, 90, 96, 101, 143, 157, 178,
 185, 194, 288, 289, 294, 311, 350, 351, 352
TLD, 44
TMS (Traffic Management Specification), 209
TNR, 161
Token-bucket, 313
Token-Ring, 86
Topologie, 49
TOS (*Type of Service*), 271, 294, 301
Traceroute, 351
Traduction d'adresses, 113
Traffic shaping, 205, 209, 291
Trame Ethernet, 123
Transceiver, 52, 103
Trunk (Ethernet), 102, 230
TTL
 ARP, 136
 DNS, 392
 IP, 286, 350, 351, 364

U

UDP, 136
Unicast, 133, 136, 259, 307
Uplink (port/lien), 55, 91, 92, 98, 227
URL, 33
UTP (câble), 72

V

V.24, 148
V.35, 149, 190
V.90, 13
VDSL, 174
Vidéo, 76, 248
VLAN, 93, 157, 225, 228
VoIP, 319, 326, 335
 VoIP, VoFR, VoATM, 255
Volumétrie, 180
VPI (Virtual Path Identifier)/VCI (*Virtual Channel Identifier*), 211
VPN/VPN-IP, 167

W

WAN, 8, 119, 236, 303, 324
WAN (*Wide Area Network*), 140
Web, 17, 26, 30, 42, 140, 144
WFQ, 295, 312, 339
WFQ (*Weighted Fair Queueing*), 290
Winipcfg, 371
WINS, 137, 262
Winsock, 136, 262, 310
WRED, 291, 295, 312
www, 27, 42, 46, 380, 390, 395, 396

X

X.121, 223
X.500, 39
X21/V11, 148, 171, 190

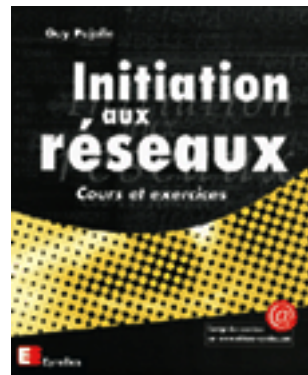
Table des encarts

À quoi sert l'adressage ?	61	Paire torsadée	73
Adressage IP (RFC 791)	111	Partage de données (T.120)	346
Adresses MAC (IEEE 802.3)	125	PIM-SM (RFC 2362)	278
Analyseur réseau	354	Pourquoi un plan d'adressage ?	110
ARP (RFC 826)	134	PPP (RFC 1661, 1662)	152
ATM (ITU I.361)	214	Proxy ARP (RFC 1027)	155
BOOTP (RFC 951, 1542)	371	Qu'est ce qu'un commutateur ?	89
Câbles cuivre en paires torsadées	53	Qu'est ce qu'un modem ?	13
Canal de contrôle H.245	332	Qu'est ce qu'un pont ?	147
Champ TOS (RFC 791, 1349)	295	Qu'est ce qu'un protocole de routage ?	238
Chantier de câblage	83	Qu'est ce qu'un réseau ?	12
Comment fonctionne un commutateur ?	98	Qu'est ce qu'un réseau local ?	48
Commutateurs Ethernet	226	Qu'est ce qu'un routeur ?	112
Commutation Frame-Relay (ITU Q.922)	200	Qu'est ce qu'un segment ?	55
Compatibilité électromagnétique	75	Qu'est ce qu'une adresse de messagerie ?	35
Composants d'un système de câblage	70	Qu'est ce qu'une adresse IP ?	16
DHCP (RFC 2131)	368	Qu'est ce qu'une carte Ethernet ?	52
DHCP options (RFC 2132)	364	Qu'est ce que FTP ?	32
Diffserv (RFC 2474, 2475)	297	Qu'est ce que HTTP et HTML ?	28
DNS (Base de données)	387	Qu'est ce que la signalisation ?	196
DNS (RFC 1034, 1035, 1995, 1996, 2181)	383	Qu'est ce que le nommage Internet ?	26
Drivers, NDIS, ODI	58	Qu'est ce que le RTC ?	14
DVMRP (RFC 1075)	267	Qu'est ce que le spanning tree ?	105
Encapsulation dans ATM	213	Qu'est ce que SNMP ?	55
Encapsulation dans Frame-Relay	198	Qualité de service Ethernet (802.1p)	299
Encapsulation IP dans Ethernet	123	Que veut dire couches réseau ?	15
Ethernet (IEEE 802.3)	91	Quels cordons de brassage ?	94
Ethernet et Token-Ring	49	Réseaux Ethernet	86
Ethernet full duplex 802.3x	91	Réseaux locaux, étendus et intersites	164
Fibre optique	77	RNIS (ITU série I)	158
Files d'attente (FIFO, WFQ, RED)	292	RSVP (RFC 2205 à 2210)	308
Files d'attente (rôle)	290	RTCP(RFC 1889 et 1890)	342
Frame Relay (ITU Q.922 XE "Q.922")	195	RTP(RFC 1889 et 1890)	341
ICMP (RFC 792, 950, 1256)	352	Signalisation d'appel Q.931	322
IGMP (RFC 1112, 2236)	262	Signalisation Frame-Relay	199
Infrastructure d'un système de câblage	68	Signalisation ILMI 4.0 (ATM Forum)	217
Interface DXI (ATM Forum)	212	Signalisation RAS (H.225)	330
IntServ (RFC 1633)	306	Signalisation UNI 4.0 (ATM Forum)	210
Inverse ARP (RFC 1293)	202	SNMP v1 (RFC 1157, 2571, 2572)	358
IPv4 (RFC 791)	133	Spanning tree (IEEE 802.1d)	127
Locaux techniques	69	TCP (RFC 793)	137
Mais qu'est ce que TCP/IP ?	16	UDP (RFC 768)	136
Messagerie sur Internet ?	35	URL	27
MIB (RFC 1155, 1212, 1213, 2863)	361	VLAN (IEEE 802.1q)	229
MOSPF (RFC 1584)	271	VRRP	235
OSPF (RFC 2328)	240		



Retrouvez nos eBooks sur :
www.ebooks.eyrolles.com

Également disponibles :



Distribution numérique par
www.GiantChair.com